

pourlascience.com

■ POUR LA

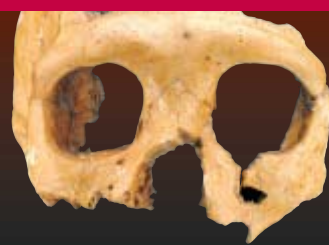
SCIENCE

Septembre 2001

Édition française de Scientific American

La formation de l'Atlantique Nord

Nos ancêtres cannibales



Les biofilms



L'univers de l'antiquité



Canada 5\$ / USA 7\$ / Royaume-Uni 3.99\$ / France 3.99\$

12687-087-30.00



BLOC-NOTES

de Didier Nordon



SCIENCE ET ÉCONOMIE

Pitié pour les troisièmes

de Ivar Ekeland

TRIBUNE DES LECTEURS



POINT DE VUE

Plus fort que l'astrologie la gènologie!

de André Lannanay



SCIENCE ET GASTRONOMIE

Huile d'olive et santé

de Hervé This

PRÉSENCE DE L'HISTOIRE

Un mathématicien mort sous la torture

de Gérard Tronel

PERSPECTIVES SCIENTIFIQUES



■ Narines des dinosaures ■ Attention, l'ombre va avoir un accident! ■ Tendre calcaire et vieilles gravures ■ Stalagmites sous-marines ■ Les visions de la pythie ■ Trous noirs et sursauts gamma ■ Une mite exceptionnelle ■ La biodiversité ■ La panique modélisée ■ Une sexualité de survie ■ L'homme à face plate ■ Cerveau et respiration ■ Les aurores à profonds ■ Un trou dans la banquise



LOGIQUE ET CALCUL

L'enfer des paris

de Jean-Paul Delahaye

ÉNIGMATH

Entrez dans le quadrille

de Dennis Shasha



ART ET SCIENCE

Les moirés de Vasarely

de Loïc Mangin



DÉES DE PHYSIQUE

La plongée sous-marine

de Roland Lehoucq, Jean-Michel Courtès



ANALYSES DE LIVRES

■ Lamarck, genèse et enjeux du transformisme 1770-1830, de Pietro Corsi ■ Le Soleil, le génome et Internet. Les révolutions scientifiques et leurs instruments, de Freeman Dyson ■ Introduction à la psychologie du sport, de Karine Bui-Xuan ■ Rencontres entre artistes et mathématiciennes. Toutes un peu les autres, de T. Chotteau, F. Delmer, P. Jakubowski, S. Pavcha, J. Peiffer, Y. Perrin, V. Roca, B. Taquet

Nos ancêtres, les cannibales

de Tim White

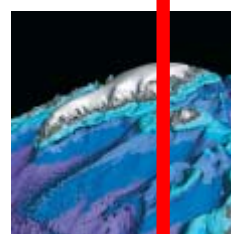
Les preuves de cannibalisme sont rares, mais il est désormais avéré que cette pratique est ancrée dans notre histoire. On en a découvert des preuves en Europe, notamment en France, et en Amérique.



La naissance de l'Atlantique Nord

de A. Braun et G. Marquart

Longtemps, la dorsale qui sillonne l'Atlantique a hésité, avant de s'orienter vers le Nord-Ouest de l'Islande. L'altimétrie satellitaire et la tomographie sismique nous révèlent comment.



Les biofilms

de B. Costerton et P. Stewart

Des bactéries organisées en films résistent aux antibiotiques et aux défenses immunitaires. Les microbiologistes étudient ces communautés bactériennes pour trouver des moyens de lutte plus efficaces.



L'Univers de Georges Lemaître

de Dominique Lambert

Georges Lemaître fut l'un des fondateurs de la théorie du Big Bang. Nombre de ses intuitions, qu'il défendit parfois contre Einstein lui-même, se sont révélées d'une portée fondamentale.



Encarts d'abonnement entre les pages 16 et 18, encart brochure service lecteurs et carte d'abonnement entre les pages 96 et 97.
En couverture : A. Braun ; données : Smith et Sandwell, GTOPO30, IBCAO ; reconstitution 3-d : Deutsche Fernstudien- und Datenzentrum - DLFB Fernstudien.

Histoire démographique et génétique du Québec

par Alain Gagnon, Hélène Vézina et Bernard Brais

Grâce aux registres paroissiaux anciens, les démographes ont reconstitué la généalogie des Québécois dont les ancêtres étaient venus de France.



7

La lumière ralentie

par Lene Verstergaard Hau

Des expériences où la lumière est ralentie, voire même arrêtée, ouvrent la voie, notamment, à de nouvelles méthodes de communication optique et à la simulation de trous noirs en laboratoire.

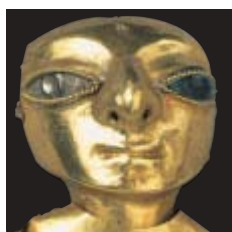


7

Or et céramiques des Andes du Nord

par Jean-François Bouchard

Depuis quelques années, les archéologues ont enfin accès à la nécropole équatorienne de Tumbaco-La Tolita, et reconstituent l'histoire des Amérindiens qui y ont prospéré vers le début de notre ère.



7

Le son des timbales

par A. Chaigne, P. Joly et L. Rhaout

Quels phénomènes physiques déterminent la qualité musicale d'un son? Pour le savoir, les acousticiens élaborent de nouvelles techniques de simulation des vibrations de l'instrument et de l'air.



7



Chaque mois, retrouvez le sommaire complet de la revue **en ligne** avec pour chaque article une bibliographie et un complément d'information.

www.pourlascience.com

Georges Lemaître, créateur de l'Univers

Les questions ont parfumé les pensées de notre enfance : comment flottent les astres de la voûte étoilée? Qu'existait-il avant que la Terre existe? Et, évidemment, qui nous a créés? Mathématiciens, physiciens et cosmologistes ont précisé ces questions... Avant l'abbé Georges Lemaître, les scientifiques observaient et s'interrogeaient sur les changements dans l'Univers ; avec Lemaître, l'Univers, en tant qu'entité, est devenu objet d'études (voir *L'Univers de Georges Lemaître*, pages 54 à 61). Est-il plus notable discontinuité épistémologique?

L'explosion des connaissances en physique dans les années 1920 a fourni les outils – mécanique quantique, relativité générale – pour cette remontée dans le temps. Il n'empêche : il fallait une énorme confiance dans les lois de la physique pour penser qu'elles s'appliquaient en des temps si éloignés (ils ne pouvaient l'être plus). L'Univers est ainsi passé d'un ensemble vague à un objet d'études caractérisé par des propriétés, objet sur lequel il était possible de raisonner physiquement. Lemaître émet l'idée que l'Univers est en expansion et que, pour connaître ses premiers instants, il suffit de dérouler l'histoire à l'envers. De surcroît, raisonne Lemaître, les traces de passé ne sont jamais totalement effacées et il doit exister des reliques de la création, quand l'Univers était l'atome primitif qu'il avait imaginé. Notre bon abbé s'est trompé dans l'identification des rayonnements fossiles témoins des premiers instants, mais il importe peu : des erreurs, permettent – la démarche est habituelle – prédictions et vérifications expérimentales, qui cernent de plus près la vérité.

Le Big Bang n'est pas une invention américaine, mais l'idée lumineuse d'un prêtre belge qui a aussi ouvert les réflexions sur le futur. Si l'Univers est né, il a un avenir sur lequel nous nous interrogeons encore, selon les voies tracées par Lemaître. Les questions de l'enfance n'ont pas toutes été résolues, mais depuis 1927, nous examinons notre Univers comme une vivante personnalité.

PHILIPPE BOULANGER



BLOC-NOTES

DIDIER NORDON



La dynastie des Ing

Mon living donnant sur un dancing, je dors peu. Je compense grâce au training autogène. Sitôt levé, mon timing se déroule selon un planning précis. D'abord, footing, jogging et body-building pour être super-top lors de mon prochain trekking. Ensuite, un rapide brushing, car je dois me présenter smart à ma holding, spécialisée dans le consulting international et le lobbying. Qu'importe si le parking est plein : je fais du rolling. Monté dans le building, j'écoute le briefing du managing, participe à un brain-storming avec l'ingéniering, surveille les mailings, vérifie les listings... Un vrai zapping! Cependant, le time-sharing informatique n'est pas pour moi. Que voulez-vous, je ne suis pas né dans le cocooning, je ne suis pas adepte des strings et autres piercings. Autrement dit, je date du baby-booming : le surging, je n'ai pas le feeling pour.

Je déplore ce perpétuel surbooking. Hélas, ce qui me manque est indigne de mon standing. J'hésite à le dire. Allons, soyons courageux! Voici mon outing. Oui, je l'avoue, j'aimerais m'accorder, par-ci par-là, un petit siesting.



Recrutement impossible

Charlatan est, en soi, une profession. Une profession utile, puisqu'il s'agit d'aller dans les roires arracher les dents pourries qu'ils voulaient bien s'y présenter. Comme l'arrachage est toujours pénible, un charlatan digne de ce nom doit savoir prononcer les phrases qui promettent monts et merveilles au patient.

Qu'est-ce alors qu'un charlatan au sein de la corporation des charlatans (charlatan au carré)? Pour le dire autrement, quand déclarera-t-on qu'un charlatan est un charlatan? Est-ce quand, pratiquant la profession divinement bien, il sait non seulement vendre du vent, mais encore en vendre alors qu'il n'y en a pas. Bref parvien

à convaincre le patient que telle dent, saine et réelle, doit être arrachée? Est-ce au contraire quand, indigne de sa profession, c'est-à-dire scrupuleux jusqu'à l'excès, il n'ose pas arracher une dent pourtant menaçante?

Le recrutement objectif des charlatans se révèle impossible. Boni de Castellane ne pensait pas autrement : il n'imaginait pas faire partie d'un club qui aurait l'incongruité de l'accepter pour membre.

La coïncidence fait le hasard

Vous marchez sur les Champs-Élysées. Dans la foule qui déambule, vous croisez Dupont. Vous ne le connaissez pas, il ne vous connaît pas, tout est banal, rien ne se passe.

Vous marchez sur les Champs-Élysées. Dans la foule, vous croisez Martin. Vous vous jetez dans ses bras : «Toi ici! Si je m'attendais à ça!» Depuis 20 ans que nous ne nous étions vus, quel hasard extraordinaire!

Etrange. La probabilité de croiser Dupont et celle de croiser Martin sont *a priori* égales. Or, la première rencontre passe inaperçue et la seconde semble marquée par l'intervention inopinée du hasard, et surprend. La notion de hasard est subjective. Il n'y a de hasard qu'à partir du moment où un homme s'étonne.

Un théorème de Bourbaki

Avoir des idées est plus facile que les mettre en ordre. Un auteur pourra passer des mois entiers à intervenir et réintervenir les chapitres au sein de son livre, les paragraphes au sein des chapitres, les phrases au sein des paragraphes, les mots au sein des phrases – voire les lettres au sein des mots! Parfois, la catastrophe naît d'un geste infime : il restructure un bref passage et voilà soudain que tout l'ensemble est, une fois de plus, à reconstruire. Inévitablement, vient un moment où l'auteur se désespère. Et si ce processus était divergent? Et si cette effrayante combinatoire était destinée à ne jamais connaître de fin?

Non. Un heureux théorème de Bourbaki affirme que le processus de réarrangement converge toujours vers un état stable : il s'arrête nécessairement au bout d'un temps fini. À vrai dire, Bourbaki n'a pas démontré son théorème, mais il l'a utilisé copieusement et avec succès, ce qui revient au même. Comme on sait, Bourbaki était un auteur collectif. Or le processus de réarrangement, qui a déjà grande

peine à satisfaire un auteur unique, n'a plus aucune chance de converger s'il y a plusieurs auteurs. Jamais ils ne se mettront d'accord. Alors, quand les bourbakistes en avaient par-dessus la tête de s'entre-injurier, quand ils ne supportaient plus de voir le dix, vingt ou trentième projet de rédaction du même volume, ils arrêtaient le jeu, envoyaient le tout à l'éditeur dans l'état où ça voulait bien se trouver et considéraient que c'était parfait comme ça.

Que les auteurs en proie aux affres de la mise en ordre de leurs idées se rassurent. Leur supplice connaîtra une fin. Théorème : toute écriture de texte converge au sens de Bourbaki.

On note les correcteurs

Lorsque les candidats à un concours sont nombreux, leurs copies sont réparties entre plusieurs correcteurs. La pratique pose problème : comment éviter qu'un candidat soit avantagé ou lésé selon le correcteur sur lequel il tombe? Avec la foi dans les statistiques, qui fait sa force, l'Administration a un credo simple : que les moyennes des notes mises par les divers correcteurs soient à peu près égales et aucune injustice n'aura été commise.

Du coup, chaque correcteur vit dans l'attente d'être repéré comme trop généreux, ou trop sévère, par rapport à ses collègues. Car il est alors soumis à la procédure dite d'«harmonisation», qu'on pourrait appeler de façon moins élégante, mais plus réaliste, «arrêter d'éviter de faire des vagues» : on lui demande de reprendre ses corrections pour essayer de se rapprocher de la moyenne générale. Tâche supplémentaire et combien ennuyeuse!

Moralité. Les candidats rêvent d'être au-dessus de la moyenne. Leurs examinateurs, plus modestes, rêvent d'être pile dedans.





Pitié pour les troisièmes!

VAR EKELAND

Les marchands offrent des options qui tirent part des caractéristiques secrètes de leurs clients, lesquels en profitent aussi. Tout est bien.

On n'est pas dans la tête d'autrui. Vous voyez ce que je fais, vous écoutez ce que je dis, vous lisez ce que j'écris, mais vous ne savez pas ce que je pense. Cette constatation paraît d'une affligeante banalité, propre tout au plus à alimenter des réflexions métaphysiques sur le sujet pensant. Il n'est rien : elle a des conséquences sur l'activité humaine, et les économistes modernes ont le mérite de l'avoir enfin mise au centre de la théorie.

Examinons un exemple simple. Voici que je m'établis comme assureur automobile. Au terme d'une étude statistique poussée, je constate que les mauvais conducteurs ont plus d'accidents que les bons, et j'établis donc deux tarifs, l'un pour les bons, l'autre pour les mauvais, calculés de manière à me laisser un modeste bénéfice. Voici donc que le premier client se présente, et je lui pose donc la question : « J'ai deux tarifs, l'un pour les bons conducteurs et l'autre pour les mauvais. Êtes-vous bon ou mauvais conducteur ? »

Notez bien que dans beaucoup de pays, où les compagnies d'assurances sont en concurrence, il leur est interdit de se communiquer les dossiers des clients. La situation que je décris, si elle n'est pas applicable en France, où les gens traînent leur *bonus* ou leur *malus* avec eux, n'en est pas moins réaliste. Qui dans ces conditions, répondrait qu'il est mauvais conducteur et réclamerait le tarif élevé ? Aussi devrais-je donc assurer tous mes clients au prix que j'avais calculé pour les bons conducteurs ; je perdrais de l'argent et je devrais relever les primes l'année suivante, amenant ainsi le départ de quelques bons conducteurs qui ne veulent pas payer pour les mauvais. Je me retrouverais avec une clientèle composée uniquement de mauvais conducteurs.

Supposons maintenant qu'au lieu d'interroger directement le client sur ses capacités de conducteur, je lui pose la question suivante : « J'ai deux tarifs : l'un avec une franchise élevée et une prime basse, l'autre avec une franchise basse et une prime élevée. Lequel choisissez-vous ? » Si ces montants sont bien calculés, les clients se répartiront en deux populations. Si les montants des primes sont bien

calculés, les bons conducteurs, ou ceux qui se jugent tels, choisiront le premier contrat (prime basse), et les mauvais le second (franchise basse). Ce calcul est délicat, car il faut que chacun des contrats soit avantageux pour la population à laquelle il est destiné, sans le devenir pour l'autre, mais enfin il est possible. Le résultat est que chacun, dans son propre intérêt, me dira la vérité, les bons conducteurs qu'ils sont bons, ce qui n'a rien d'extraordinaire, et les mauvais qu'ils sont mauvais, ce qui est plus remarquable, non par des mots, mais par le contrat choisi.

Si l'ensemble des deux contrats que l'assureur propose permet d'atteindre ce résultat, on dit que ces contrats sont révélateurs (ou incitatifs), et l'assureur détermine alors, parmi tous les contrats révélateurs, le plus avantageux pour lui (qui lui rapportera le plus d'argent). Ses possibilités de gain ne sont pas illimitées car, s'il matraque trop les clients, ils iront s'assurer ailleurs et il y a donc un meilleur contrat révélateur. Celui-ci a des propriétés intéressantes.

On constate, par exemple, qu'il faut faire des conditions particulièrement avantageuses aux bons conducteurs pour qu'ils ne soient pas attirés par le contrat proposé aux mauvais. En d'autres termes, si l'on pouvait lire sur les visages qui conduisent bien et qui conduisent mal, l'assureur changerait ses contrats d'une façon qui n'attirerait pas les mauvais conducteurs mais où les bons conducteurs seraient assurés à un « juste prix » moins avantageux pour eux que celui dont ils bénéficient en l'absence d'information. En d'autres termes, ceux-ci, du fait qu'ils sont seuls à savoir qu'ils sont bons et qu'ils peuvent prétendre qu'ils sont mauvais, touchent une rente dite « informationnelle ».

On pourrait penser que ce genre de considérations est typique du néolibéralisme actuel et du souci d'extraire des poches des consommateurs jusqu'au dernier centime, mais l'on se tromperait. Les racines de cette théorie se trouvent en France, et même dans le débat sur le service public à la française, entamé depuis deux siècles et toujours d'actualité.

La première étude sur les contrats révélateurs est le mémoire que Léon Walras adresse en 1875 à Jules Ferry, alors ministre des Transports, comme contribution au débat sur la nationalisation des chemins de fer. Dans ce mémoire, Walras, qui devait passer à la postérité comme le fondateur de la théorie de l'équilibre général, explique la politique tarifaire des compagnies de chemin de fer et, notamment, l'existence des trois classes pour les voyageurs, à la lumière de ce qui n'est pas encore la théorie des contrats révélateurs. En fait, dit-il, dans un français admirable : « Les compagnies considèrent, à tort ou à raison, le prix moyen de 7,06 centimes (au kilomètre), prix assez voisin de celui de 7,5 qui est le prix des secondes, comme étant le prix de bénéfice maximum ; mais elles ne veulent pas négliger d'accepter plus des voyageurs disposés à payer plus, ni même refuser d'accepter moins des voyageurs décidés à ne pas payer tant... »

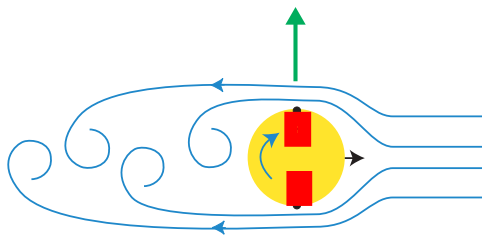
Quand Walras écrit, les compagnies sont privées, mais cela ne change rien : dans la logique du Service public, l'Etat propriétaire chercherait à ce que les voyageurs qui peuvent payer plus subventionnent ceux qui ne peuvent pas payer tant. Ne pouvant connaître la déclaration de revenus des voyageurs lorsqu'ils retiennent leur place, la SNCF atteint le même résultat en leur proposant aussi des premières ou des secondes. Laissons pour terminer la parole à Walras :

« De la l'institution des trois classes distinctes et les grands efforts qu'elles font pour accentuer, d'une part les avantages de la 1^{re} classe, et, d'autre part, les désavantages de la 3^e. Quand on réclamait jadis si grands cris la fermeture des voitures de 3^e classe par des vitres, quand on réclame aujourd'hui leur chauffage en hiver, et qu'on se plaint à ce propos de la dureté des compagnies, on ne saisit pas leur vrai mobile. Si les voitures de 3^e classe étaient assez confortables pour que beaucoup de voyageurs de 2^e classe et quelques-uns de 1^{re} y allaient, le produit net total, tel qu'il se compose d'après la théorie du monopole, serait abaissé. Et voilà tout ! »

TRIBUNE DES LECTEURS

Je la «liffe» ou je la «coupe»?

Il me semble que l'article de Mickael Dickinson *Le vol des insectes* (voir *Pour la Science*, août 2001) comporte une erreur. Un joueur de tennis qui «coupe une balle» il lance la balle en avant et la fait tourner en même temps dans le sens inverse des aiguilles d'une montre) ne voit pas sa balle s'élever, tel que le montre la figure 7 de l'article, mais descendre. Pour la voir s'élever, le joueur doit la «liffer», c'est-à-dire la lancer en la faisant tourner dans le sens des aiguilles d'une montre. La raison en est l'effet Magnus. Durant sa course, la balle est suivie d'un sillage turbulent constitué de tourbillons tournant alternativement dans un sens et dans un autre (tourbillons de Kármán, voir la figure ci-dessous). Quand la balle est en rotation dans le sens indiqué sur la figure, la vitesse relative de la balle par rapport à l'air (somme de la vitesse de l'air et de la vitesse tangentielle de la balle) est plus grande au point *H* qu'au point *B*. Or, d'après le théorème de Bernoulli, la pression d'un fluide est d'autant plus petite que la vitesse du fluide est grande. La pression



en *H* est donc inférieure à la pression en *B* et cette différence crée une force de portance dirigée vers le haut (flèche verte). Cet effet est utilisé pour propulser des navires où les voiles sont remplacées par des cylindres tournants.

Roland MORISOT
Boulogne-Billancourt

Miroir, mon beau miroir dis-moi...

Dans *La naissance de la vie* (voir *Pour la Science*, juin 2001), Robert Hazen rappelle que toute molécule se trouvant dans une certaine configuration spatiale peut exister dans une autre configuration qui est l'image de la première dans un miroir. L'une des configurations est dite «gauche» et l'autre «droite». Selon R. Hazen, la prédominance dans les organismes vivants d'acides aminés «gauches» par rapport aux «droits» serait le fruit du pur hasard. Or dans l'article, on présume que la sélection des variétés gauche et droite des acides

aminés se fait via les faces de cristaux où les acides se seraient assemblés. Les faces de cristaux, elles aussi, existent dans deux configurations, images l'une de l'autre dans un miroir. Puisque les deux versions des cristaux peuplent la nature en proportions égales, pourquoi n'en est-il pas de même pour les acides aminés? Le hasard me paraît difficilement responsable. Pourquoi ne pas envisager que de chaque côté du miroir vivent deux mondes aux lois dissemblables? À un moment donné, une molécule «droite» n'aurait pas été aussi efficace à remplir une certaine tâche que son homologue gauche et la sélection naturelle, faisant son œuvre, aurait brisé le miroir.

Aurélien DUPRE
Nantes

Réponse de Robert Hazen

J'ai deux raisons pour soutenir l'hypothèse du hasard pur. D'abord, pour toute réaction chimique qui mène à une prédominance de molécules gauches devant les droites, il existe le même mécanisme de l'autre côté du miroir qui conduit au résultat opposé. Ensuite, même dans le cas où un camp l'aurait emporté devant l'autre, à un moment de l'histoire de la Terre, au cours des temps géologiques, cette hégémonie se serait nécessairement éteinte.

Arche perdue ou Saint-Graal?

Très honoré d'être cité dans *Pétra. Une capitale aux confins du désert* (voir *Pour la Science*, août 2001), je dois cependant vous signaler une erreur commise dans l'évocation de mes aventures. S'il est vrai que mes pérégrinations m'ont conduit jusqu'à l'ancienne cité de Pétra, la motivation que vous me prêtez est erronée. J'y recherche, non pas l'Arche d'Alliance comme vous l'avez écrit, ce coffre de bois qui aurait contenu les tables de la loi reçues par Moïse sur le mont Sinaï, mais le Saint-Graal, la coupe dans laquelle Jésus aurait bu le vin de sa dernière Pâques et qui aurait servi à recueillir le sang de la crucifixion.

Indiana JONES

La Terre devance la mer

Dans l'article *Les caprices des marées* (voir *Pour la Science*, août 2001), Bernard Simon écrit que la pleine mer a toujours un retard sur la culmination de la Lune, c'est-à-dire sur le moment où la hauteur de la Lune au-dessus de l'horizon est maximale, ce qui en raison de la symétrie du mouvement, doit se produire lorsque la Lune passe au méridien. Or dans l'article *Les danses spatiales*

de la Lune (voir *Pour la Science*, juillet 2001), il est dit que le bourrelet solide formé à l'équateur de la Terre par l'attraction de la Lune se produit toujours un peu en avance sur la direction Terre-Lune. Pourquoi les marées sont-elles en retard sur la Lune et le bourrelet solide en avance?

Michel FEUILLE
Franconville

Réponse de Bernard Simon

Le mécanisme des marées terrestres est très éloigné de celui des marées océaniques dans lesquelles les courants jouent un rôle essentiel. En fait, vis-à-vis de la force génératrice de la marée la Terre réagit comme un corps élastique, tandis que les océans réagissent comme un fluide où des ondes sont produites de manière diffuse en chaque point. La marée observée en un point est le résultat de la superposition de ces ondes qui ont effectué des parcours divers avec les réflexions, diffractions et dissipations associées. On comprend intuitivement, de ce fait, que la marée observée est en retard sur la force génératrice : il faut laisser le temps à ces ondes d'atteindre l'observateur. Ce retard se nomme l'âge de la marée. Dans le monde, il est en moyenne d'une quarantaine d'heures, mais varie considérablement autour de cette valeur.

Motifs damassés sur météorites

Je voudrais ajouter quelques compléments à l'article *Le mystère des épées de Damas* (voir *Pour la Science*, août 2001). Le phénomène d'exsolution dont se nourrit le métal damassé est un mécanisme général qui agit par la ségrégation des différents constituants d'un métal lors du refroidissement, produisant des motifs dans le métal. Les figures dites de Widmannstätten que l'on observe sur les sections polies de roches extraterrestres sont le résultat de la ségrégation, à la suite d'un refroidissement lent d'un minerai de fer pauvre en nickel et d'un de fer riche en nickel. L'exsolution se fait encore plus géométrique dans le cas des structures réticulées où l'un des constituants du minéral oriente la répartition de l'autre.

Gérard TEMEY
Espay Saint-Marcel

Écrivez à *Pour la Science* vos réactions, vos commentaires et vos réflexions sur l'actualité scientifique et vos relations avec la science. L'adresse de *Pour la Science* sur Internet est : tribune@pouirlascience.fr. Découvrez les autres contributions de lecteurs sur le site www.pouirlascience.com.



Plus fort que l'astrologie, la génologie!

ANDRÉ LANGANEY

Pour retrouver le passé et prédire le futur, l'examen de l'ADN est la nouvelle charlatanerie.

Un bon créneau commercial consiste à faire croire au bon peuple ravi, de manière non réfutable, ce qu'il désire le plus et que personne n'est en mesure de lui offrir dans l'état actuel des connaissances. «Grandes» religions et sectes vendent ainsi du surnaturel et de l'arbitraire social. Sorcelleries et astrologies vendent du futur rassurant ou inquiétant selon votre degré de stress ou de dépendance. Le génie des astrologues est de s'être moulés depuis l'Antiquité dans l'autorité abusive de la science pour imposer leur charabia et leur commerce : comment ne pas croire celui qui parle de manière incompréhensible avec tant d'aisance? Et quand des universitaires doctorisent une voyante, on se demande si le pire reste possible.

Il l'est! Nos mages et charlatans ne seraient pas allés plus loin que l'information des horoscopes sans l'aide des universitaires. Aussi fallait-il que l'Université vienne au secours des escrocs dépassés par les technologies. C'est fait! La science la plus médiatisée, la génétique, vient d'y contribuer par deux fois, timidement certes, mais ce n'est qu'un début. Le combat, n'en doutons pas, se poursuivra!

Il y a deux ans, un chercheur inconnu d'une université de troisième zone, aux États-Unis, s'était rendu compte que les angoisses du passé valent parfois celles de l'avenir. En particulier, pour les résilients des communautés meurtries par ce passé, comme les Noirs américains. Beaucoup voudraient savoir de quelle région d'Afrique venaient leurs ancêtres déportés, de quelle ville ou de quel village, quelle était leur religion (nous sommes aux États-Unis!), leur structure sociale (un peu de polygamie pour les plus loustics, un peu d'Islam pour les plus radicaux)... Puisqu'ils veulent le savoir, il faut leur fournir l'information pour un prix acceptable, maximisant le rapport coût-bénéfice – 300 \$ est un bon chiffre. On ouvrit donc un site Internet pour donner ces bonnes nouvelles du passé, moyennant finances. Restait à gérer un détail : l'information personnelle de base. Les lignes de la main et le

marc de café ne faisant plus sérieux, on prit l'ADN, qui apporte une quantité d'informations illimitée, récupérable pour passer, en plus ou moins bon état, sur des trottis buccaux ou sur des cheveux. Cela dit, même avec les techniques de séquençage les plus sophistiquées et les plus coûteuses, les biologistes ne peuvent circonscrire l'origine probable d'un sujet qu'à une région très étendue. De plus, la recherche du pays d'origine est quasiment impossible quand la personne appartient à un groupe aussi métissé que les Noirs américains ou aussi homogène génétiquement que la famille linguistique Nigéro-Congo, à laquelle appartenaient presque tous leurs ancêtres africains!

Vous concluez donc que le projet annoncé est irréalisable? Vous ne m'avez pas bien suivi : nous parlons commerce et l'objectif est de satisfaire le client avec une information qui ne risque pas d'être démentie, même si elle ne figure pas dans son ADN!

Qu'il y ait des escrocs au pays des gangsters, quoi de plus naturel? Ce n'est pas à Oxford, au Collège de France ou à l'Académie des sciences que des scientifiques joueraient à cela? Erreur! Dans votre supermarché vous trouverez la traduction d'un livre de Bryan Sykes, d'un prestigieux institut médical d'Oxford. Il y explique que tout Européen descend d'une seule femme parmi les sept «filles d'Ève». Selon l'auteur, ces dernières descendraient d'un seul lignage maternel venu d'Afrique et qui serait à l'origine de tous les «non-Africains» (concept discutable!). Cette «Ève non africaine» serait elle-même issue de l'hypothétique «Ève africaine» qui aurait, par descendance féminine, transmis l'ADN mitochondrial (un petit chromosome extranucléaire) à toute notre espèce. Que l'ADN mitochondrial ne soit qu'un deux cent millièmes du génome, que la théorie de l'Ève africaine ait été réfutée dans les meilleures revues spécialisées, en 1989 par Laurent Excoffier, puis, deux ans plus tard, par d'autres, ne trouble pas Bryan Sykes! Il va jusqu'à préciser les lieux et

dates exacts où vivaient nos plantureuses ancêtres, décrit leurs charmes et caractères répartis selon les stéréotypes de réputations ethniques des habitants actuels de l'endroit. Là encore, un site Internet pour 200 £ cette fois, vous permet d'envoyer votre ADN à une start-up d'une école de série d'Oxford, laquelle fournit en retour le roman de votre supposée ancêtre. La préhistoire est «reconstituée» à partir de quelques données archéologiques triviales sur les temps et lieux supposés originels de la pom-pom Ève... et de pas mal d'imagination (feuilletez le livre sans l'acheter si vous ne me croyez pas!).

Vous pensez que ce jeu n'amuse que les midinettes et les clients de la foire du Trône? Une revue souvent mieux inspirée a fait une couverture très médiatisée sur le sujet, avec photo d'un éminent du Collège de France, secrétaire perpétuel de l'Académie, posant complaisamment avec une ex-Miss France. Dans le texte, chacun retrouvait «son» Ève, mais l'illustration suggérait que le professeur descendait de la belle Ève, montrant que le progrès esthétique n'a rien d'une fatalité!

Il n'y a pas de raison de s'arrêter là et d'écouter les ennuyeux qui expliquent que la reconstitution de l'histoire des populations ancestrales par la génétique est difficile, méticuleuse, et fournit peu de réponses simples aux questions du public. Nos collègues anglo-saxons sont si «héréditaristes» depuis Francis Galton, si convaincus que les gènes déterminent tout de nous depuis Edward Wilson ou Simon LeVay, pères respectifs de la «sociobiologie» et du «gène de l'homosexualité»), que n'importe quelle séquence d'ADN sera supposée aussi prédictrice que les boules de cristal pour faire de nouveaux horoscopes et une «astrologie génétique» beaucoup plus compliquée et impressionnante que les thèmes astraux. On l'appellera génologie, pour entretenir la confusion avec la génétique, et l'avenir promis s'y révélera meilleur produit que le passé imaginé.

André LANGANEY est professeur au Musée de l'Homme et à l'Université de Genève.



Huile d'olive et santé

HERVE THIS

On a identifié dans l'huile d'olive les composés qui préviennent les méfaits des agents oxydants

Si nous ne sommes pas ce que nous mangeons (*Man ist was man isst* dit l'adage allemand), le fait est là : certaines alimentations sont plus saines que d'autres. À tort ou à raison, les aliments fumés consommés par les populations du Nord de l'Europe favorisent l'apparition de cancers du système digestif, alors que l'alimentation du bassin méditerranéen, riche en huile d'olive, en fruits, en légumes et en poissons, prévient ces cancers, ainsi que les maladies cardiovasculaires.

Ces résultats résultent de l'étude épidémiologique la plus importante jamais conduite sur les relations entre l'alimentation et le cancer : depuis 1992, des chercheurs de dix pays d'Europe collaborent à l'EPIC (*European Prospective Investigation Into Cancer and Nutrition*), une étude portant sur un demi-million de personnes suivies ! Des prises de sang régulières, stockées dans l'azote liquide pour des analyses ultérieures, les mensurations des sujets enregistrées, ainsi que leur état de santé... À Lyon, en juin dernier, les premiers résultats ont été dégagés : une consommation quotidienne de 500 grammes de fruits et de légumes diminue de moitié l'incidence des cancers des voies aérodigestives, et réduit notablement celle des cancers du côlon ou du rectum ; le tabac et l'alcool ont des effets désastreux sur les cancers des voies aérodigestives supérieures. Ainsi le risque d'un de ces cancers pour quelqu'un qui fume un paquet par jour est huit fois supérieur au risque pour un non-fumeur. Une consommation d'éthanol supérieure à 60 grammes par jour, soit environ 75 centilitres de vin, multiplie par neuf le risque d'un de ces cancers.

L'extraordinaire banque de sang qui a été constituée sert aussi à des études plus ciblées. Ainsi, à Heidelberg, Helmut Bartsch et ses collègues du Centre allemand de recherches cancérologiques DKF2 ont cherché à corréler la consommation de plusieurs huiles et des modifications de l'ADN dans les globules blancs. Ils supposaient, d'une part, que l'auto-oxydation des lipides (le « rancissement ») provoquait la formation de composés réactifs

nuisant aux cellules, et, d'autre part, que ces composés réactifs se liaient à l'ADN provoquant des mutations qui engendraient *in fine* des cancers.

Les chercheurs allemands ont reculé le sang de femmes qui participaient à l'étude EPIC et qui continuaient de se nourrir comme à l'habitude : les femmes qui avaient le moins d'ADN modifiés étaient celles qui consommaient le plus de fruits et de légumes. Ainsi le dosage des ADN modifiés est-il un bon test *in vivo* de l'efficacité des antioxydants, tels ceux identifiés dans l'huile d'olive.

Pourquoi cet intérêt pour l'huile d'olive ?

Parce que les études épidémiologiques font état d'une incidence réduite des cancers et des maladies cardiovasculaires chez les populations méditerranéennes. Depuis la découverte de cette corrélation, de nombreux facteurs ont été invoqués pour expliquer les bienfaits du « régime méditerranéen » : l'abondance des tannins dans les vins locaux, la diversité de l'alimentation, la faible consommation de produits fumés, la forte consommation de fibres, la consommation notable de fruits et de légumes, la consommation d'huile d'olive... L'huile d'olive est un « bon » produit alimentaire à plus d'un titre. Comme les autres huiles, ses molécules majoritaires, les triglycérides, sont composées d'une molécule de glycérol à laquelle se lient trois acides gras. Toutefois, plus de 70 pour cent de ses acides gras sont l'acide oléique, mono-insaturé avec une seule double liaison entre des atomes de carbone du squelette de l'acide gras, lequel est moins sensible à l'auto-oxydation que les acides gras polyinsaturés, tels que l'acide linoléique de l'huile de tournesol. Or l'oxydation des acides gras polyinsaturés engendre des molécules réactives qui se lient – on le sait par des études *in vitro* – à l'ADN de cellules humaines et qui risquent de favoriser des cancers.

D'autres propriétés de l'huile d'olive expliqueraient son intérêt nutritionnel. À Heidelberg, R. Owen a analysé diverses huiles d'olive à la recherche des composés antioxydants, et a montré que ces pro-

priétés protectrices proviennent aussi de la présence de plusieurs phénols (des molécules contenant toutes un cycle à six atomes de carbone, lié à au moins un groupe hydroxyle -OH).

Lors de la croissance des olives, entre octobre et janvier, les principaux composés phénoliques sont liés à des sucres. Mais, lors de la maturation des fruits, des enzymes séparent le sucre des phénols, qui sont libérés. De nombreux phénols ont été identifiés dans l'huile d'olive dont l'hydroxytyrosol et deux « lignanes ». Tous ont de fortes propriétés antioxydantes, et tous sont plus abondants dans les huiles extra-vierges que dans les huiles raffinées.

L'étude de l'auto-oxydation des huiles

montre que l'hydroxytyrosol est le composé le plus actif : le peroxyde d'hydrogène (présent dans l'« eau oxygénée »), qui réduit considérablement la viabilité des cellules épithéliales (cultivées) de l'intestin, est complètement inhibé par l'ajout d'hydroxytyrosol. R. Owen et ses collègues ont montré que 8 huiles d'olive extra-vierges et 5 huiles d'olive raffinées avaient ce même effet protecteur, alors que 7 huiles de tournesol ne l'avaient pas. Les composés phénoliques de l'huile d'olive sont des antioxydants puissants.

Les lignanes ont également été étudiés pour leurs effets protecteurs contre les cancers. Ils semblent agir différemment : leur structure, analogue à celle de l'estradiol, en ferait des anti-estrogènes qui bloquent la prolifération des cellules du sein). L'huile d'olive contient aussi de fortes concentrations de squalène, composé qui est transféré à la peau (le sébum en contient 12 pour cent) et qui agirait contre les cancers de la peau.

La gastronomie incitait à utiliser l'huile d'olive. La biologie confirme que bon goût et hygiène vont de pair.

Prochain rendez-vous France Info et Pour la Science, le 26 septembre, avec la chronique Info Sciences de Marie Odile Monchicourt.

■ POUR LA SCIENCE

VOTRE ABONNEMENT 100 % NUMÉRIQUE

1 an au prix de 48 €

~~soit 4 € seulement par numéro au lieu de 5,20 €~~

Pour consulter et télécharger
en toute liberté sur www.pourlascience.fr
les 12 prochains numéros
dès leur parution.



Oui, je m'abonne
à Pour la Science 100 % numérique
1 an (12 numéros) pour 48 €

Je préfère découvrir toutes
les offres d'abonnement disponibles
en cliquant ici



Un mathématicien mort sous la torture

GERARD TRONEI

Maurice Audin, adepte de Bourbaki et militant pour l'indépendance de l'Algérie, achevait sa thèse lorsque les parachutistes l'ont arrêté à Alger, en juin 1957.

Parmi les innombrables tragédies que la guerre d'Algérie a causées, l'une d'elles se détache par le retentissement qu'elle a eu dans l'opinion publique, à la fin de la période coloniale, et parce qu'elle a frappé une communauté particulière, celle des mathématiciens. C'est le drame «Maurice Audin».

Pendant l'été 1957, la presse de l'«anti-France» annonçait qu'un jeune assistant de mathématiques de la faculté d'Alger, Maurice Audin, avait été arrêté le 11 juin à son domicile, par des parachutistes, et avait été officiellement porté disparu dix jours plus tard, le 21 juin. C'est du moins la thèse que les autorités françaises ont soutenue à l'époque, et soutiennent encore aujourd'hui, plus de quarante ans après les faits. Une version officielle de sa disparition a été montée de toutes pièces. Audin se serait évadé au cours d'un transfert. En réalité, Audin a été torturé à mort. Ses assassins, promus et décorés depuis, coulent des jours paisibles ou racontent leurs «exploits», alors qu'ils mériteraient être condamnés pour apologie de crimes de guerre.

Dès l'automne 1957, un comité s'est constitué pour faire connaître l'affaire Audin, ainsi que de nombreux autres cas de tortures. Cependant, si l'on a beaucoup parlé de la mort de ce mathématicien militant, sa vie et son œuvre, dont il sera question dans cet article, sont restées méconnues. Le destin d'Audin, scellé alors qu'il n'avait que 25 ans, rappelle celui d'Évariste Galois, lui aussi fauché en pleine jeunesse (à 21 ans, lors d'un duel) au début du XIX^e siècle. Cependant, Galois avait choisi de combattre son adversaire pour des raisons personnelles, sachant ce qu'il encourait. Audin n'a fait que lutter, avec la seule force de ses convictions, contre les choix politiques de la France pendant la guerre d'Algérie.

Qui était Audin, et qu'a-t-il eu le temps d'accomplir en mathématiques? Audin naît en Tunisie, en 1932, d'un père postier. Il est le cadet d'une famille de quatre enfants. Ses parents déménagent bientôt à Alger, et, en 1942, il entre au lycée Gauthier, mais l'établissement est réquisitionné pour servir de caserne lorsque les troupes américaines débarquent en Afrique du Nord. Son père est devenu gendarme, et Audin est envoyé dans une école d'enfants de troupe, gratuite en échange de quelques années de service dans l'armée à la fin des études. Il ne s'y plaît guère, mais y reste jusqu'à ce que ses parents soient en mesure de rembourser ses frais de scolarité, c'est-à-dire jusqu'en première. De retour au lycée Gauthier, il obtient son baccalauréat de mathématiques haut la main et s'inscrit à la Faculté des sciences d'Alger. Pour subvenir à ses besoins, il donne des cours particuliers, comme de nombreux étudiants issus de milieux modestes.

En licence, il rencontre le professeur René De Possel, qui enseigne l'analyse et la mécanique. De Possel est l'un des fondateurs du groupe Bourbaki. Ce collectif de mathématiciens, créé en 1934, a modernisé la discipline, clarifiant et réorganisant son contenu. Son style concis et rigoureux a fait de nombreux adeptes. Bourbaki a privilégié la méthode axiomatique. Autrement dit, on énonce clairement et précisément les définitions, les règles de bases (les axiomes) auxquelles doivent obéir les entités mathématiques considérées et l'on en déduit, par des chaînes de raisonnement strictement logiques, d'autres propriétés moins évidentes (les théorèmes). Entre 1950 et 1980 en particulier, Bourbaki a changé la face des mathématiques françaises.

De Possel remarque immédiatement Audin, qui est le meilleur de sa promotion. En même temps qu'Audin termine

le certificat de physique générale, bête noire des étudiants en mathématiques, mais nécessaire à l'obtention de la licence de mathématiques, il commence ses recherches et publie une première note aux *Comptes-rendus* de l'Académie des sciences, en 1953. La même année, il est nommé assistant de mathématiques à la Faculté des sciences d'Alger, à 21 ans, peu de temps après son mariage avec Josette Audin, dont il a eu trois enfants (Michèle, Louis et Pierre), n'a pas cessé de se battre pour connaître la vérité sur la mort de son mari. Elle vient de déposer une plainte pour crimes contre l'humanité et séquestration.

LES UNIVERSITAIRES ENTRENT EN GUERRE

En dehors des mathématiques, Audin a une seconde activité : la politique. Il milite d'abord avec les étudiants communistes puis au Parti communiste algérien. Il combat le colonialisme et participe à la lutte politique pour libérer l'Algérie. Souvenons-nous qu'après les massacres de Sétif, le 8 mai 1945, les Algériens s'organisent progressivement contre la France. En 1954, ils déclenchent une insurrection armée et créent le Front de libération nationale (FLN), qui rassemble les mouvements nationalistes algériens. En réaction, les autorités françaises interdisent les partis nationalistes, puis le parti communiste algérien. Les activités politiques d'Audin deviennent clandestines, mais il ne participe pas directement à la lutte armée.

C'est dans cette situation difficile qu'il assure ses cours, suit les séminaires organisés par De Possel et les dernières avancées de la recherche en mathématiques. Il lit beaucoup, en particulier les *Éléments de mathématique* de Bourbaki, les notes aux *Comptes-rendus* de l'Académie des sciences, les *Doklady* (les Comptes

endus de l'Académie des sciences soviétique, en langue russe, qu'il a appris suffisamment pour déchiffrer les articles (qui l'intéressent).

En octobre 1956, ses recherches sont assez avancées pour qu'il envisage de soutenir sa thèse l'année suivante. Il vient à Paris pour s'inscrire à la Faculté des sciences et demander à Laurent Schwartz de faire partie du jury. Ce dernier est membre du groupe Bourbaki ; il a reçu la médaille Fields en 1950 (l'équivalent du prix Nobel, absent en mathématiques pour sa théorie des distributions. Les distributions sont un élargissement du concept de fonction de Dirac, important dans l'étude des équations aux dérivées partielles. L. Schwartz relève quelques erreurs mineures dans le travail du jeune mathématicien. La thèse n'est pas tout à fait prête, mais Audin fait part à son aîné de son désir de la terminer sans tarder, car il craint une détérioration des conditions politiques en Algérie, où les militaires poursuivent les opposants avec acharnement. À son retour à Alger, Audin se consacre à la rédaction de sa thèse qu'il a presque finie lorsqu'il est arrêté. Le lendemain, Henri Alleg, qu'il avait rencontré au Parti communiste, est lui aussi arrêté et torturé. Dans son livre intitulé *La Question*, H. Alleg raconte avoir croisé son ami après le 11 juin : celui-ci portait déjà de nombreuses marques de torture. Après cette rencontre, la trace d'Audin disparaît, et l'affaire Audin commence.

À la rentrée 1957, De Possel propose à Laurent Schwartz de faire passer la thèse du disparu *in absentia*. La thèse est soutenue en décembre, dans une salle comble, en présence de nombreux journalistes. Elle relance le combat des universitaires contre les tortures en Algérie.

LES RECHERCHES MATHÉMATIQUES EN 1957

L'œuvre mathématique d'Audin, naturellement émouvante, illustre les préoccupations des jeunes mathématiciens de l'époque. Poursuivie entre 1953 et 1957, elle est contenue dans quatre notes aux *Comptes-rendus* de l'Académie des sciences, qui jalonnent l'avancée de sa thèse, et du manuscrit de la thèse elle-même, intitulée *Sur les équations linéaires dans un espace vectoriel*.

Le style du manuscrit est proche de celui de Bourbaki, qui exerce à l'époque une grande influence sur les mathématiciens français. Audin est inspiré par Bourbaki à travers De Possel. Ce dernier étudiait toutes sortes de problèmes, qu'ils soient abstraits ou concrets, telles les oscillations dites de lacet des wagons de chemin de fer sur leurs rails. Cet éclectisme suggère que le professeur n'au-



Pendant la guerre d'Algérie, la presse qui portaient à la connaissance du public les tortures et les disparitions était souvent saisie. *L'Humanité* du 7 mars 1961 est parue avec une page blanche, marquée en son centre du mot «censuré» (en haut à droite). L'article qui devait paraître portait sur «l'affaire des harkis»; il donnait la parole à deux Algériens, torturés dans une cave du 18^e arrondissement de Paris, comme beaucoup d'autres (en haut à gauche). À l'époque, une commission de censure officielle décidait de ce qui pouvait être publié. Quand le journal était prêt au tirage, l'employé de la censure venait réclamer un double des preuves originales. Si le journal était saisi, on imprimait une nouvelle édition en catastrophe avec des «blancs». Le premier finissait au pilon. La pratique de la censure rejoint l'attitude des généraux et du pouvoir civil dans l'affaire Audin : ils ont accumulé les faux témoignages pour dissimuler l'assassinat du jeune mathématicien (en bas).





ments, appelés vecteurs, peuvent s'additionner entre eux et être multipliés par des «nombres» appelés scalaires : si $\mathbf{v}$ et $\mathbf{w}$ sont deux vecteurs, $\mathbf{v} + \mathbf{w} = \mathbf{w} + \mathbf{v}$ est encore un vecteur, et si α est un scalaire, $\alpha\mathbf{v}$ est aussi un vecteur. L'un des exemples les plus simples d'espace vectoriel est l'ensemble des vecteurs du plan, qui s'additionnent selon la règle du parallélogramme et qui peuvent être multipliés par des nombres réels. On dit qu'il s'agit là d'un espace vectoriel «sur $\mathbb{R}$ », les scalaires étant des nombres réels. On peut aussi construire des espaces vectoriels sur $\mathbb{C}$ (les scalaires sont des nombres complexes) ou plus généralement, sur d'autres ensembles possédant certaines structures abstraites, que l'on appelle des «corps».

Un espace vectoriel topologique est un espace vectoriel où l'on peut passer continûment d'un vecteur à l'autre (la topologie se préoccupe de tout ce qui dépend de la notion de limite, de continuité et de voisinage).

Plus précisément, c'est un espace vectoriel où l'application qui, à tout couple de vecteurs $(\mathbf{v}, \mathbf{w})$, fait correspondre le vecteur $\mathbf{v} + \mathbf{w}$, et l'application qui, à tout scalaire α et tout vecteur $\mathbf{v}$, fait correspondre le vecteur $\alpha\mathbf{v}$, sont continues. La théorie des espaces vectoriels topologiques généralise la théorie des espaces vectoriels normés (espaces où, pour chaque vecteur, on peut définir une «norme», une longueur en quelque sorte), tels les espaces vectoriels euclidiens. Dans les espaces vectoriels topologiques, il n'y a pas nécessairement de notion de distance, mais il y a toujours une notion de continuité. Cette théorie est utile en analyse fonctionnelle, soit l'étude des espaces de dimensions infinies dont les éléments sont des fonctions.

Dès lors, on perçoit pourquoi Bourbaki et ses disciples se sont penchés sur ces espaces vectoriels topologiques, qui englobent les espaces plus classiques : les propriétés des équations linéaires que l'on y démontre restent vraies dans des espaces plus petits, de sorte que l'on n'a pas besoin de recommencer la démonstration. Attaché à la concision et à la précision, Bourbaki voulait présenter les mathématiques de la façon la plus générale possible. Il voulait fournir à ses lecteurs des outils mathématiques «aussi robustes et aussi universels que possibles» selon les mots d'André Weil, un autre de ses fondateurs. Le «paquet abstrait» de l'œuvre de Bourbaki, à savoir les notions générales qu'il voulait inclure au début et qui seraient utiles pour l'ensemble des sujets abordés ensuite, a grossi de plus en plus à mesure que le travail avançait. Dans ce paquet, les mathématiques apparaissent comme un édifice doté d'une unité profonde, reposant notamment sur la théorie des ensembles, et hiérarchisé en termes de structures abstraites (algébriques, topologiques...).

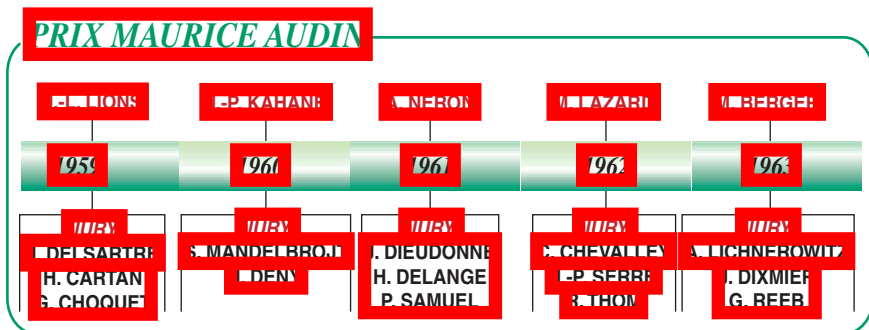
L'un des fils conducteurs de la thèse d'Audin est l'étude des propriétés d'opérateurs linéaires dans les espaces vectoriels topologiques de dimensions infinies. Un opérateur linéaire, notons le A , est tel que $A(\alpha\mathbf{v}_1 + \beta\mathbf{v}_2) = \alpha A(\mathbf{v}_1) + \beta A(\mathbf{v}_2)$, quels que soient les vecteurs $\mathbf{v}_1$ et $\mathbf{v}_2$ et les nombres scalaires α et β .

Parmi les sujets à la mode dans les années 1950 figure la généralisation de propriétés des opérateurs linéaires bien connues dans les espaces de dimensions finies (où les vecteurs sont caractérisés par un nombre fini de composantes), aux espaces de dimensions infinies. En dimensions finies, de nombreuses propriétés topologiques des opérateurs linéaires se déduisent de leurs propriétés algébriques, ce qui, en général, est faux en dimensions infinies. Ainsi, quand on considère les solutions x et y d'un système de deux équations à deux inconnues, soit $a_{11}x + a_{12}y = b_1$ et $a_{21}x + a_{22}y = b_2$, on montre facilement que $x + \varepsilon$ et $y + \varepsilon$ sont des solutions (voies de x et y) du système où a_{11} est remplacé par $a_{11} + \varepsilon$, a_{12} par $a_{12} + \varepsilon$, puisque $a_{11}\varepsilon + a_{12}\varepsilon$ et $a_{21}\varepsilon + a_{22}\varepsilon$ sont négligeables. En revanche, en dimensions infinies, c'est à-dire pour un système d'une infinité d'équations à une infinité de dimensions, on ne pourra plus négliger les sommes $a_{ij}\varepsilon$, car une somme infinie d'éléments infiniment petits n'est pas forcément infiniment petite. Audin voulait ainsi contribuer à franchir un pas de plus dans certaines démarches de généralisation entreprise par Bourbaki, en passant des espaces à dimensions finies à des espaces à dimensions infinies.

Il n'aurait pas imposé un sujet à son étudiant et qu'Audin a lui-même choisi un thème «bourbachique». D'ailleurs, la bibliographie de la thèse d'Audin confirme son intérêt pour les volumes des *Éléments de mathématique* sur les espaces vectoriels topologiques, parus dans les années 1950.

LES ESPACES VECTORIELS TOPOLOGIQUES

Quelles sont les caractéristiques de ces espaces et pourquoi Bourbaki, puis notre jeune mathématicien, les ont-ils étudiés ? La structure d'espace vectoriel est fondamentale en algèbre linéaire et en analyse, vastes domaines issus de l'étude des systèmes d'équation linéaires, aux très nombreuses applications. Un espace vectoriel est un ensemble dont les éléments



Créé en 1959, le prix Audin a récompensé chaque année un jeune mathématicien pendant la guerre d'Algérie. Il a été décerné une dernière fois en 1964, à Paul-André Meyer et Pierre Cartier. Un prix Audin a également été décerné en Italie en 1961.

DES MATHÉMATIQUES UNIVERSELLES

Il a également exploré plusieurs autres directions de recherche dans son entre-prise généralisatrice. Il voulait démontrer, dans les espaces topologiques abstraits, des propriétés établies dans des espaces vectoriels particuliers (par exemple, des propriétés des opérateurs intégraux dans des espaces de fonctions intégrables).

Sa thèse contient aussi des notions fondamentales sur des propriétés spectrales de familles d'opérateurs linéaires. La théorie spectrale est une extension de l'étude du problème dit «des valeurs propres et des vecteurs propres» dans le cas des espaces vectoriels de dimensions infinies, tels que certains espaces de fonctions (par exemple, l'espace des fonctions $f: \mathbb{R} \rightarrow \mathbb{C}$ à valeurs complexes telles que l'intégrale du carré $|f|^2$ de leur module, donne un nombre fini ; ces fonctions jouent un rôle important, non seulement en mathématiques, mais aussi en physique et en chimie quantiques). Le problème des valeurs propres et des vecteurs propres est relatif aux applications linéaires opérant sur un espace vectoriel V . Si A désigne une application linéaire de cet espace, et s'il existe un vecteur non nul v et un nombre scalaire α tels que $A(v) = \alpha v$, v est dit vecteur propre et α est appelé valeur propre de A associé à v .

Examinons un exemple dans le plan : si P est la projection sur la droite OX parallèlement à OY , on peut voir que tout vecteur v parallèle à OX reste inchangé sous l'action de P ; autrement dit, on a $P(v) = v$; la valeur propre est 1, et comme vecteur propre associé on peut choisir n'importe quel vecteur dirigé selon OX . Mais si l'on considère un vecteur v parallèle à OY , on voit que $P(v) = 0 = 0.v$: la valeur propre est 0, et comme vecteur propre associé, on peut choisir n'importe quel vecteur de direction OY . Les nombres 0 et 1 sont les seules valeurs propres de cette projection P . L'ensemble des valeurs propres d'une application linéaire A s'appelle le spectre de A en dimensions finies et fait partie du spectre de A en dimensions infinies. L'intérêt d'étudier les vecteurs et les valeurs propres d'une application linéaire est de caractériser ses propriétés et de simplifier certains calculs.

Dans la dernière partie de sa thèse Audin commence à explorer quelques directions de recherche qu'il pensait probablement approfondir par la suite. Il traite des équations linéaires par la méthode de Schmidt, qui consiste à décomposer des opérateurs linéaires à dimensions infinies en somme d'opé-

rateurs plus simples. Cette méthode déjà utilisée dans le cadre hilbertien soit dans les espaces vectoriels euclidiens à dimension infinie, commençait tout juste à être employée dans le cas d'espaces plus généraux, tels les espaces de Banach (des espaces où il n'y a pas de notion d'orthogonalité) et des espaces vectoriels topologiques. Audin entame l'étude des perturbations des opérateurs que l'on introduit dans des équations linéaires et qui modifient légèrement ou qui «perturbent» les solutions) et leurs effets sur les propriétés topologiques de familles d'opérateurs linéaires. Il projetait sans doute aussi de démontrer certaines conjectures que le mathématicien polonais Stephane Banach, l'un des pères de l'analyse fonctionnelle, avait énoncé en 1932 dans sa *Théorie des opérateurs linéaires*, mentionné dans la bibliographie de sa thèse. Il semble avoir été influencé par les *Leçons d'analyse fonctionnelle*, des mathématiciens hongrois Francis Riesz et Bela Nagi. Dans sa thèse, il présente quelques applications à des espaces particuliers de fonctions de propriétés démontrées dans des espaces plus abstraits.

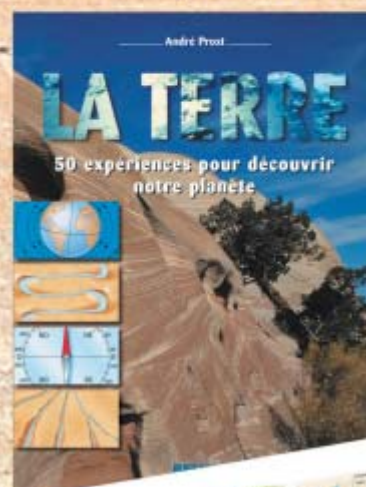
Au début de l'affaire Audin, comment la communauté mathématique a-t-elle réagi ? Une certaine presse, souvent clandestine et souvent saisie, contenait des informations sur la guerre d'Algérie et son cortège de tortures, de disparitions et de mensonges des militaires couverts par le gouvernement. De même, les articles de François Mauriac, à l'*Express*, de Pierre Henri Simon au *Monde*, les livres de l'historien Pierre Vidal-Naquet, les témoignages de Henri Alleg, devenu journaliste à l'*Humanité*, et de militaires de carrière, tel le général Paris de la Bollardière, ou d'apôlles portaient la situation à la connaissance du public. Cependant, la plupart des contemporains de Maurice Audin, les mathématiciens et les autres, ne voulaient pas savoir. La soutenance de la thèse d'Audin a eu un retentissement bien vite étouffé, mais pour maintenir son souvenir vivace, un prix Audin a été créé en 1959. Pendant la guerre d'Algérie, il a récompensé chaque année un jeune mathématicien.

Si l'on n'évoque pas périodiquement les événements douloureux du passé, la mémoire s'estompe. Pourquoi ne pas recréer un prix Audin qui distinguerait un jeune mathématicien de 25 ans, soit l'âge d'Audin lorsqu'il a disparu ?

Gérard TRONEL a fait sa carrière de mathématicien à l'Université Pierre et Marie Curie.

LA TERRE

50 expériences pour découvrir notre planète



Comment naissent les dunes ?
Comment une rivière creuse-t-elle une vallée ?
Comment se forment les grottes ?
Pourquoi les pierres éclatent-elles sous le froid ?
Ce livre propose une série d'expériences pour comprendre comment le vent, la pluie, la glace, les cours d'eau modelent le visage de la Terre.

ANDRÉ PROST est géologue, ancien enseignant à l'université d'Orléans et à Paris VI, et ingénieur au Bureau des Recherches Géologiques et Minières à Orléans.

Code 002401-100 F-128 pages

Voir bon de commande page 96

BELIN • POUR LA SCIENCE

Les narines des dinosaures

La découverte de la position des narines charnues des dinosaures renseigne sur leur physiologie.

La question semble futile. Où étaient placées les narines charnues des dinosaures ? Pourtant, la réponse est d'importance, car elle renseigne sur la façon dont ces animaux respiraient et maintenaient leur température interne. Localiser les narines charnues signifie les situer par rapport à l'ouverture des fosses nasales du crâne : les narines sont dites caudales, quand elles sont à l'arrière des fosses et rostrales, à l'avant. Chez les dinosaures, on a longtemps cru qu'elles étaient caudales. Or, Lawrence Witmer, de l'Université de l'Ohio, a montré qu'elles étaient rostrales.

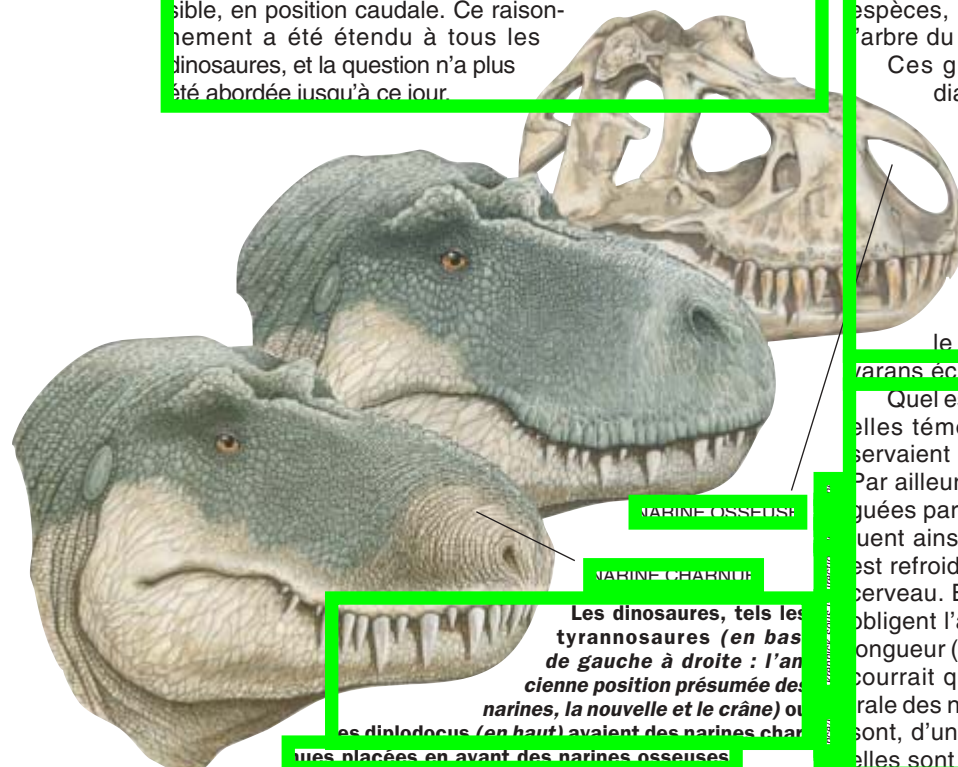
Chez les êtres humains, la taille réduite des fosses nasales osseuses laisse peu d'ambiguïté sur la place des narines charnues. Cependant, chez les dinosaures, et particulièrement chez les sauropodes (des dinosaures à quatre pattes et à long cou), les ankylosaures (des animaux à bec de canard) et les cératopsidés (leur crâne est surmonté de cornes), les fosses nasales osseuses mesurent jusqu'à 60 centimètres, soit près de la moitié de la longueur du crâne. Il est alors difficile, en l'absence d'animaux vivants, de situer les narines charnues.

L'idée de sauropodes amphibies a longtemps prévalu (à la fin des années 1950, on a cependant montré que ce n'était pas le cas) et a influencé les paléontologues. Selon eux, ces animaux utilisaient leurs narines à la façon d'un tuba pour respirer : leurs narines devaient donc être le plus haut possible, en position caudale. Ce raisonnement a été étendu à tous les dinosaures, et la question n'a plus été abordée jusqu'à ce jour.



Par radiographie, L. Witmer a comparé la position des narines charnues à celle des narines osseuses chez des espèces, encore vivantes, voisines des dinosaures dans l'arbre du règne animal : des oiseaux et des crocodiliens. Ces groupes d'animaux composent le groupe des diapsides, c'est-à-dire ceux qui ont deux paires de fosses temporales, en arrière de l'orbite oculaire. Afin que les narines charnues soient visibles, il les a enduites au préalable de sulfate de baryum, qui est opaque aux rayons X. Chez tous les animaux observés, les narines charnues sont en position rostrale. De ces observations, on déduit que les narines rostrales étaient la norme chez les dinosaures. C'est aussi le cas des mammifères et des tortues. Seuls les varans échappent à cette règle.

Quel est l'avantage de telles narines ? Placées en avant, elles témoignent d'une cavité nasale volumineuse, qui servaient à réchauffer, à humidifier et à filtrer l'air inspiré. Par ailleurs, les parois charnues des narines sont très irriguées par de nombreux vaisseaux sanguins : elles constituent ainsi une région d'échanges thermiques où le sang est refroidi par contact avec l'air inspiré avant d'arriver au cerveau. Enfin, des narines charnues placées vers l'avant obligent l'air à traverser les cavités nasales sur toute leur longueur (avec des narines charnues caudales, l'air ne pourrait que la partie supérieure). Enfin, la position rostrale des narines est plus adaptée à l'odorat, car les narines sont, d'une part, davantage exposées à l'air, d'autre part, elles sont proches de la bouche et des aliments ingérés.



NARINE OSSEUSE

NARINE CHARNE

Les dinosaures, tels les tyrannosaures (en bas de gauche à droite : l'ancienne position présumée des narines, la nouvelle et le crâne) ou les diplodocus (en haut) avaient des narines charnues placées en avant des narines osseuses.

Attention, l'ombre va avoir un accident!

Les enfants âgés de moins de huit ans prennent les ombres comme des objets à part entière.

Lucky Luke est plus rapide que son ombre : voilà qui ne doit pas étonner outre mesure les enfants âgés de moins de huit ans. Michèle

Molina et François Jouen, de l'université de Rouen et de l'École pratique des hautes études à Paris, viennent de réaliser une expérience qui montre que les enfants considèrent les ombres comme des objets solides et compacts, distincts des objets qui les créent.

De nombreux psychologues ont essayé de comprendre comment les enfants perçoivent et conçoivent les objets physiques. D'après les études réalisées par Elisabeth Spelke, de l'université de Cornell, les enfants distinguent les objets dans l'environnement grâce à l'utilisation de trois principes fondamentaux : le principe de contact (l'action sur un objet n'est possible que s'il y a contact) ; celui de cohésion (la forme de l'objet ne varie pas quand il se déplace) et celui de continuité (les trajectoires sont continues). La contrainte de solidité des objets découle de ces trois principes. En revanche, rares sont les études sur ces étranges traces sombres qui accompagnent les objets : les ombres.

Le psychologue suisse Jean Piaget avait déjà défini plusieurs stades de conception des ombres chez les enfants. Au début, ils considèrent que l'ombre correspond à un objet isolé. Puis, à partir de l'âge de huit ans, ils établissent une relation entre l'ombre et une source lumineuse, mais ils croient toujours que les ombres continuent à exister la nuit. Vers l'âge de neuf ans seulement, les enfants comprennent le rôle de la lumière dans la production des ombres. Avant neuf ans, ils considèrent l'ombre

comme une entité, mais lui attribuent-ils les mêmes propriétés qu'aux objets?

Pour répondre à cette question, M. Molina et F. Jouen ont réalisé une expérience auprès de 30 enfants, classés en trois groupes selon leur âge : quatre ans, six ans, et huit à neuf ans. Ils cherchaient à déterminer si les enfants appliquent aux ombres (comme aux objets) la contrainte de solidité, bien que les ombres violent les trois principes liés aux objets. Dans l'expérience, une voiture miniature se déplace sur un rail. Une lampe l'éclaire et projette son ombre sur un écran derrière la voiture. Un obstacle (un tas de galets) est placé sur la trajectoire de l'ombre. On demande aux enfants d'imaginer ce qu'il va se passer lorsque la voiture avance et que son ombre arrive sur l'obstacle. Ils ont le choix entre trois possibilités, représentées par des dessins proposés selon un ordre aléatoire : l'ombre de la voiture grimpe sur les galets ; elle s'arrête devant ou elle passe à la surface des galets.

Les réponses des enfants âgés de quatre et six ans sont identiques. Selon eux, neuf sur dix d'entre eux, l'ombre de la voiture va s'arrêter devant le tas de galets ou alors elle va le gravir, comme une vraie voiture. Au contraire, dans le groupe des enfants de huit à neuf ans, neuf sur dix donnent la bonne réponse, contraire à la contrainte de solidité des objets : l'ombre glisse sur les galets.

Dans une seconde expérience, l'ombre de la voiture rencontre non plus le tas de galets réel, mais une autre ombre. Pour les enfants âgés de moins de huit ans, l'ombre se comporte face à une autre ombre comme un objet face à un autre objet. Seuls les plus âgés savent que les deux ombres se fondent. Ces deux expériences renforcent l'idée que les savoirs des enfants se «différencient» avec l'âge : que les enfants créent de nouvelles «catégories» à mesure de leur développement. Les psychologues vont maintenant étudier si les enfants attribuent des propriétés «humaines» aux ombres formées par des personnes.

BRÈVES

L'âge du capitaine

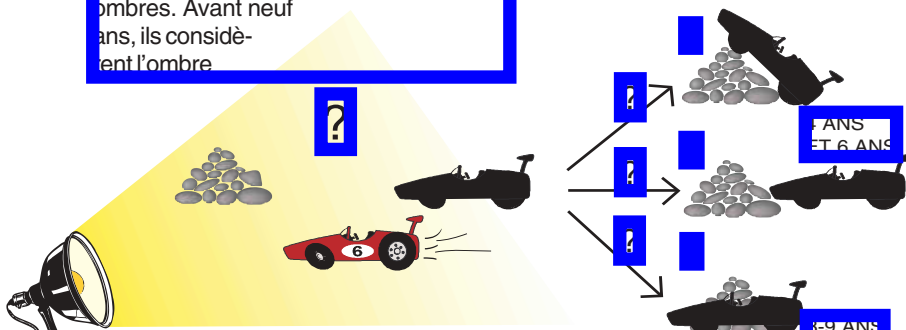
Quel âge a votre capitaine? Le capitaine étant un gros poisson des eaux douces d'Afrique, la sclérochronologie peut apporter la réponse. Cette technique consiste à compter les stries de croissance des pièces calcifiées des poissons, telles les otolithes. Le biologiste est aujourd'hui aidé dans cette tâche par un logiciel mis au point à l'Institut de recherche et de développement, l'IRD. Cet outil informatique fournit des images numériques sur lesquelles les stries de croissance sont représentées par des régions translucides alternant avec des zones opaques. La détermination de l'âge de certains poissons pourrait améliorer la gestion des ressources halieutiques.

Naine blanche et matière noire

Plus de 90 pour cent de la matière de l'Univers n'est pas détectable par la lumière qu'elle émet et ne se manifeste que par ses effets gravitationnels. On sait que, dans notre galaxie, une grande partie de cette matière est localisée dans le halo, une sphère de gaz et d'étoiles. Les astronomes pensaient qu'une partie de la matière noire s'y trouvait sous forme d'étoiles peu lumineuses, notamment des naines blanches, des restes très denses d'étoiles de type solaire. Toutefois, Céline Reylé, Annie Robin, de l'Observatoire de Besançon et Michel Crézé, de l'Université de Bretagne Sud, viennent d'établir un modèle de notre galaxie, et ont montré qu'un grand nombre de naines blanches que l'on croyait localisées dans le halo appartiennent, en fait, au disque de la Voie lactée. Les naines blanches ne semblent pas constituer plus d'un pour cent de la mystérieuse masse sombre du halo galactique. La question de l'origine de la masse noire reste ouverte.

Hibernatus avait des puces!

Le corps gelé de cet homme, surnommé l'Homme de Hauslabjoch, Ötzi ou Hibernatus, a été découvert dans les glaciers des Alpes italiennes en 1991. Il vivait il y a 5 300 ans à la frontière entre l'Autriche et de l'Italie. Une équipe de paléontologues de l'Université d'Innsbruck l'ont rapidement décongelé, puis recongelé en 1999, pour prélever des échantillons de peau, d'ADN, ou encore de vêtements. Deux puces de la même espèce, parasite de l'homme, ont été retrouvées sur ces vêtements. L'analyse ne dit pas si elles étaient savantes.



Trois groupes d'enfants respectivement âgés de quatre ans, six ans et huit à neuf ans sont interrogés : comment se «comportera» l'ombre de la voiture rouge projetée sur un écran quand elle arrivera sur l'obstacle, le tas de pierres. Selon les enfants de quatre et six ans, l'ombre passe par-dessus l'obstacle (a) ou elle s'arrête devant (b) : pour les jeunes enfants, l'ombre est un «objet». Les enfants plus âgés donnent la bonne réponse : l'ombre glisse sur l'obstacle (c).

Obésité et diabète

De nombreux obèses développent un diabète, cependant, on ignorait le lien entre ces deux maladies. Les équipes de Philippe Froguel, de l'Institut Pasteur de Lille, et de Takashi Kadowaki, de l'Université de Tokyo, ont mis en évidence le rôle d'une hormone, l'adiponectine, sécrétée par le tissu adipeux, dans l'apparition d'un diabète chez les obèses. En effet, quand le gène de cette hormone est muté, d'une part, les graisses s'accumulent dans la ceinture abdominale, d'autre part, l'organisme résiste à l'action de l'insuline, qui normalise la concentration en sucre dans le sang. De nouveaux médicaments devraient être fondés sur ces résultats.

Mélodie au piano

Comment une mélodie se dégage-t-elle de toutes les notes jouées par un pianiste ? Inconsciemment, il joue les notes de la mélodie en attaquant légèrement plus fortement les touches du piano, révèle une étude menée par Werner Goebel, de l'Institut de recherche en intelligence artificielle, à Vienne. Ce dernier a équipé les touches d'un piano de capteurs, et il a calculé la vitesse des marteaux, pour deux morceaux joués par 22 pianistes. Il a montré que les notes de la mélodie sont jouées 20 à 30 millièmes de seconde en avance par rapport aux autres notes.

Baby-blues avant l'accouchement

La dépression post-natale est reconnue, mais elle serait moins fréquente que celle qui touche les femmes avant l'accouchement. Le baby-blues est fréquent avant l'accouchement, d'après une étude menée par une équipe de psychiatres et d'obstétriciens de Bristol, sur 14 000 femmes ayant accouché : 13,5 pour cent des femmes présentent de signes de dépression après 32 semaines de grossesse contre neuf pour cent huit semaines après l'accouchement.

Méditer nuit à la santé

La fréquentation des temples et des églises pourrait se révéler dangereuse pour la santé. D'après une étude faite par des chimistes de l'Université de Cheng Kung à Taiwan, la combustion de l'encens dégage des hydrocarbures aromatiques polycycliques, qui sont des substances cancérigènes : dans un temple bouddhiste, ils ont mesuré que la concentration en benzopyrène était presque 45 fois supérieure à celle qui règne au domicile des fumeurs. Méditez, mais aérez!

Tendre calcaire et vieilles gravures

Ornée de nombreuses gravures, la grotte de Cussac livre de

Cussac est le « Lascaux de la gravure » et Lascaux « la Chapelle Sixtine de la peinture pariétale ». Ces citations, du conservateur de la grotte de Lascaux, Jean-Michel Clottes et de l'Abbé Breuil, le pionnier de l'étude de l'art pariétal, résument en quelques mots l'intérêt et l'importance de la grotte de Cussac, en Dordogne.

Lorsqu'en octobre 2000, le spéléologue Marc Delluc emmène Norbert Aujoulat, du Centre national de la préhistoire à Périgueux, et son confrère Christian Archambeau, reconnaître une grotte, aucun des deux ne s'attend à une découverte de cet intérêt. L'entrée, fort petite, est située à trois kilomètres de la rivière Dordogne. Après avoir rampé le long d'un étroit tunnel, puis s'être faufilés sous un éboulis instable, les préhistoriens parviennent dans une grande galerie de 15 mètres de largeur sur dix mètres de hauteur, apparemment très longue. Là, plus d'une centaine de gravures, à l'évidence très anciennes et d'un style nouveau, les attendent. Tout le bestiaire de l'art pariétal y est représenté : bisons, mammoths, bouquetins, chevaux, rhinocéros. À cela s'ajoutent des représentations animales rares (oiseaux, bêtes au mufler allongé,...) ainsi que des images de femmes et de triangles cubiens.

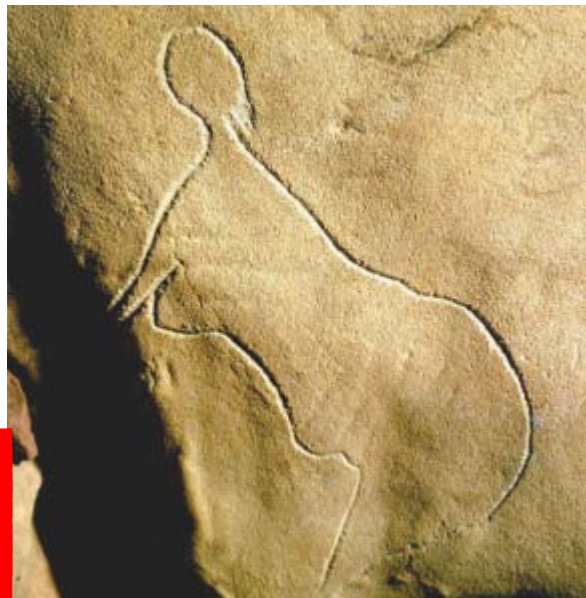
Quoi que primitives, ces gravures sont très esthétiques. Elles sont inscrites dans un calcaire ocre, assez tendre pour que les anciens artistes l'aient incisé à l'aide d'une pointe de bois ou de silex, voire avec des ongles. Elles sont d'une taille inhabituelle dans l'art pariétal : les grands panneaux combinent des figures dont certaines mesurent plusieurs mètres de hauteur. Elles ont été réalisées sur des pans rocheux mis à nu par la chute d'énormes blocs. Les surfaces lisses et planes ainsi formées s'élèvent en surplomb. Sans doute les artistes paléolithiques se juchaient-ils sur le bloc effondré afin d'atteindre les parois les plus élevées du panneau.

Les préhistoriens associent le caractère archaïque des figures et certaines des conventions graphiques qu'elles contiennent à une période ancienne du Paléolithique. Ils estiment qu'il s'agit probablement du Gravettien, une période allant de -28 000 à -22 000 ans (cette période doit son nom à celui de Gravette, situé à seulement quatre kilomètres). Le Gravet

ien a précédé le dernier maximum glaciaire qui a eu lieu vers -20 000 ans.

Selon N. Aujoulat, les formes des gravures, les thèmes et les associations thématiques permettent de rapprocher la grotte de Cussac de celle de Pech Merle dans le département du Lot. Découverte en 1922, cette grotte contient de nombreuses gravures et peintures très bien conservées, dont certaines ont été réalisées vers -25 000 ans. Les points communs entre les styles des deux grottes intéressent les préhistoriens, car les grottes de Cussac et de Pech Merle se trouvent dans la même zone interfluviale, qui sépare le Lot et la Dordogne. Toutefois, la comparaison approfondie des styles artistiques des deux grottes ne se fera qu'après l'évaluation scientifique de celle de Cussac. Pour le moment, le site a seulement été protégé (dans des conditions exemplaires grâce à son découvreur M. Delluc). Les préhistoriens vont y faire installer une plateforme de circulation en inox suspendue au niveau du sol. Cette précaution prise, ils ne pénétreront plus dans la grotte qu'en chaussons quand ils iront répertorier les œuvres et étudier les traces de passage dans la grotte. Celles-ci comprennent notamment les restes de sept squelettes d'adultes, des traces de pas, des griffures d'ours, mais apparemment aucun objet ou autre mobilier préhistorique.

Presque intact, l'un des squelettes semble être celui d'un homme mort couché sur le ventre. Très incomplets et sans connexions anatomiques, les autres ossements sont disposés en tas dans des cuvettes qui semblent avoir été des bauges d'ours. Les archéologues ont prélevé un peu de matière sur chaque squelette afin de les dater. Les résultats permettront de confirmer que ces ossements datent bien du Gravettien, ou d'une période postérieure.



Une des gravures de la grotte de Cussac

Stalagmites sous-marines

Un nouveau type de sources hydrothermales a été découvert dans l'Atlantique.

Au milieu des océans, la croûte océanique se sépare en deux plaques qui s'éloignent l'une de l'autre, laissant remonter du magma des entrailles de la Terre. Ce magma, lors de sa remontée, se refroidit et se solidifie, donnant naissance à de vastes roches, les dorsales océaniques, qui sont le siège d'une activité hydrothermale : l'eau de mer pénètre dans la croûte à travers des réseaux de fissures et descend à quelques kilomètres de profondeur où la roche est chaude. Là, au contact des roches, elle se charge en minéraux, qui, quand l'eau remonte à la surface de la dorsale, précipitent et forment par accréation des structures solides aux formes variées : les sources hydrothermales. Les plus spectaculaires sont des cheminées de plusieurs dizaines de mètres de hauteur, sortes de vestiges d'une cité engloutie. Une équipe composée d'océanographes et d'un minéralogiste vient de découvrir un nouveau type de sources hydrothermales qui se distinguent de celles déjà répertoriées par leur distance à la dorsale médio-atlantique (les sources ne sont pas situées sur la dorsale même, mais à une quinzaine de kilomètres) et par un environnement chimique original.

C'est au large du Sud du Maroc, que l'équipe menée par Deborah Kelley de l'Université de l'État de Washington a découvert ce nouveau type de sources hydrothermales. Descendus sur le plancher océanique, à 700 mètres de profondeur, à bord du sous-marin d'exploration *Alvin*, ils ont été les premiers à contempler un ensemble de structures qui se détachent du fond par leur blancheur. La « Cité perdue » regroupe des sources actives ou éteintes : flèches, monticules, cheminées (dont l'une de 60 mètres de hauteur, est la plus haute répertoriée à ce jour). Les sources présentent une arborisation importante sur leurs flancs, et une couche blanche les recouvre (signe de jeunesse de la couche minérale).

La majorité des sources hydrothermales connues naissent sur le sol jeune des dorsales océaniques composé de magma fraîchement solidifié, et émettent des fluides hydrothermaux chauffés à plus de 300 °C. Les sources de la Cité perdue, à 15 kilomètres de la dorsale, laissent échapper un fluide dont la température est inférieure : entre 40 et 75 °C.

Sur une dorsale océanique, la croûte nouvellement formée repousse la plaque océanique de part et d'autre. La vitesse d'écartement de la plaque océanique sur laquelle repose la Cité perdue est faible (1,2 centimètre par an comparée à celle de la plaque pacifique (15 centimètres par an au niveau de l'île de Pâques). Les sources hydrothermales situées en dehors des dorsales et découvertes précédemment sont localisées dans des zones qui se déplacent à des vitesses intermédiaires entre ces deux valeurs. Or, cette vitesse détermine le type de roches présentes dans la croûte. Ainsi, pour des vitesses élevées, seules des roches issues de la fusion des couches superficielles du manteau terrestre (des basaltes) constituent la croûte, tandis que pour des vitesses inférieures, par exemple celles de la Cité perdue, des roches magmatiques dont la composition est proche des roches du manteau (des péridotites) affleurent. Quand un fluide hydrothermal s'échappe d'un sol basaltique, il est très acide. Au contraire, dans le cas de la Cité perdue, le sol est riche en péridotites et le fluide est basique.

A cause de la température peu élevée et du caractère basique des minéraux qui précipitent sont des carbonates et des hydroxydes. La couleur blanche de ces minéraux contraste avec

celles des autres sources de l'Atlantique, où le fluide hydrothermal, chaud et acide, riche en fer et en composés du soufre s'échappe des cheminées sous forme de grandes fumées noires.

La découverte, à la fin des années 1970, des premières sources hydrothermales se doublait d'une trouvaille importante : la vie foisonne dans ces régions plongées dans l'obscurité. Une multitude d'organismes tirent leur énergie de réactions chimiques qui ont lieu dans cet environnement, a priori hostile du fait de l'obscurité, des températures élevées et de la présence de composés toxiques dans les fluides émis. D. Kelley et ses collègues ont découvert que la Cité perdue abrite aussi des communautés microbiennes. Ces dernières sont combinées en filaments de plusieurs centimètres de longueur ou plaquées à la surface des cheminées. Des analyses du patrimoine génétique de ces micro-organismes préciseront s'ils sont spécifiques du site.

La question a longtemps été débattue de savoir comment s'étaient assemblées à la surface terrestre les premières molécules du vivant malgré l'intense rayonnement ultraviolet qui frappait la planète dans sa prime jeunesse. L'une des hypothèses est aujourd'hui que la vie serait apparue dans l'océan, près des sources hydrothermales. À des profondeurs certaines abritées du rayonnement ultraviolet, mais privées d'énergie lumineuse, les sources représentent des oasis où les réactions chimiques qui ont enfanté la vie auraient été possibles. En outre, les sources sont riches en certains minéraux qui facilitent la formation de molécules biologiques (en servant notamment, sur les surfaces cristallines, de support pour la constitution de longues molécules à partir d'atomes solés ou bien de catalyseurs de réactions chimiques). La Cité perdue, avec ses conditions favorables (fluide basique, présence de composés carbonés, température peu élevée), constitue l'un des environnements où la vie a pu apparaître.



Structures découvertes à la Cité perdue (en haut une cheminée en bas des cristallisations désordonnées). Elles sont recouvertes d'une couche minérale blanche.

Capture en plein vol

Les grandes noctules, des chauves-souris vivant en Europe, volant jusqu'à 500 mètres d'altitude, se nourrissent exclusivement d'insectes. Du moins le croyait-on. Une équipe espagnole a montré qu'elles appréciaient aussi les petits oiseaux migrateurs nocturnes. L'analyse de leurs excréments a révélé de nombreux restes d'oiseaux passériformes, immense groupe de petits oiseaux auxquels appartiennent notamment les moineaux, et qui survolent par millions l'Espagne durant leur migration. Les noctules ont une envergure qui peut atteindre 60 centimètres, elles ne pèsent pas plus de 50 grammes : ces conditions semblent leur permettre de capturer en vol les oiseaux dont elles se nourrissent. On croyait la capture des proies en vol l'apanage de certains rapaces, mais les noctules semblent l'utiliser.

Un anticorps contre le VIH

Un anticorps humain, nommé b₁₂, inhiberait l'activité du virus responsable du SIDA. Cet anticorps a été découvert dans la moelle osseuse d'un homme de 31 ans, séropositif, mais qui ne manifeste aucun symptôme de la maladie depuis six ans. Cet anticorps se lierait à un récepteur de surface du virus et empêcherait d'interagir avec ses cellules cibles.

Echec et maths

Apprendre en jouant n'est pas l'apanage des jeunes enfants. Dès la prochaine rentrée, les étudiants de l'Université d'Aberdeen, en Grande-Bretagne, vont pouvoir suivre des cours de doctorat en science des échecs. Pendant leur thèse, les étudiants, tous joueurs d'échecs, vont tenter de créer de nouveaux logiciels de jeu d'échecs, et de transposer les modes de raisonnement des joueurs à des programmes d'intelligence artificielle.

Trio en symbiose

Drôle de cohabitation dans les eaux profondes autour de l'île d'Elbe ! Les biologistes de l'Institut Max Planck, à Brême, en Allemagne, ont découvert des organismes qui vivent en symbiose non pas à deux, mais à trois : un ver et deux bactéries. Le ver abrite les deux bactéries, qui, en échange, nourrissent leur hôte dépourvu de tube digestif. La première transforme les sulfates de l'eau de mer en sulfures, que la seconde oxyde. Une première !

Les visions de la Pythie

On aurait découvert l'origine des gaz qu'inhalait la Pythie pour rendre les oracles.

En 480 avant notre ère, les Athéniens, en guerre contre les Perses, s'enquirent auprès de l'oracle

de Delphes de la conduite à adopter devant l'offensive ennemie, mais ne reçurent en réponse que des promesses de malheur : « A la flamme furieuse, Arès (le dieu de la guerre) livrera bien des temples, où les images des immortels se dressent aujourd'hui couvertes de sueur tremblantes d'effroi, et du haut des toits ruisselle un sang noir, présage du désastre fatal. » On venait de tout le pays consulter Apollon, le fils de Zeus, par l'intermédiaire de l'oracle de Delphes. L'écrivain romain Plutarque a rapporté comment la cérémonie de l'oracle se déroulait : la Pythie, celle par qui s'exprimaient les dieux, se retirait dans une pièce située sous le temple d'Apollon où, en inhalant des vapeurs, elle entraînait en transe. Elle rendait ensuite l'oracle en des propos incompréhensibles, que seuls des prêtres savaient interpréter. Selon Plutarque, les gaz inhalés par la Pythie provenaient d'une fissure dans le sol ou d'une source d'eau. Cependant, l'excavation du temple à la fin du XIX^e siècle ne révéla ni fissure, ni trace de source. Une équipe constituée d'un géologue, d'un géochimiste, d'un archéologue et d'un médecin, propose aujourd'hui, après analyse de la géologie du sol, que celui-ci est suffisamment perméable pour laisser passer des gaz. En outre, ils ont découvert dans les eaux qui s'écoulent à proximité du site des traces d'éthylène qu'on sait responsable d'hallucinations. La diffusion de ce gaz par les roches vers l'intérieur du temple aurait été responsable des transes de la Pythie.

Delphes est localisée le long du golfe de Corinthe, dans une zone de forte activité sismique. Le golfe est le fruit de la tectonique des plaques qui repousse la partie Nord de la Grèce, l'Attique, de la partie Sud, le Péloponnèse. De nombreuses failles, parallèles au golfe, découpent le sol de la région. L'une d'elles, la faille de Delphes, passe précisément sous le sanctuaire ; de larges zones en sont visibles en surface. Jelle Deboer et ses collègues du Département des sciences de la Terre et de l'environnement de l'Université de Midd-



Zeus consultant l'oracle de Delphes

etown ont découvert une autre faille, la « faille de Kerna », perpendiculaire à la faille de Delphes, qui passe aussi sous le temple d'Apollon. À cause de la rencontre des deux failles, le sol y est riche en fissures où certains gaz, vaporisés par la chaleur, pourraient remonter en surface. On ne connaît pas le lieu exact de l'intersection des deux failles, mais, par extrapolation, elle semble située à la verticale du temple.

Y a-t-il des gaz dans le sol de Delphes susceptibles de donner des hallucinations ? Les gaz viendraient directement des roches ou seraient véhiculés sous forme dissoute par une source d'eau. On connaît aujourd'hui à proximité du temple deux sources, mais des recherches géologiques et archéologiques montrent qu'il devait en exister, dans l'Antiquité, au moins huit : à plusieurs endroits du sanctuaire, on trouve des travertins, des dépôts de minéraux laissés par une source. L'équipe a analysé l'eau qui s'écoule en contrebas du temple, ainsi que des échantillons des travertins, qui portaient peut-être la trace des gaz dissous dans les anciennes sources. Ils ont trouvé de l'éthylène dans la source et, dans les travertins, des traces de méthane et d'éthane. Le méthane et l'éthane sont faiblement toxiques, mais en s'accumulant dans la pièce où se retirait la Pythie, ils pouvaient produire des hallucinations. Quant à l'éthylène, c'est une molécule anciennement utilisée comme anesthésique. À faible dose, elle donne une légère sensation d'euphorie. À doses plus élevées, elle provoque délire et inconscience. Ces constatations concordent avec ce que Plutarque relate sur les vapeurs prophétiques : elles avaient une odeur sucrée et elles ont entraîné la mort d'une Pythie qui en avait trop inhalé.

Trous noirs et sursauts gamma

Les sursauts gamma sont les explosions les plus violentes de l'Univers. La formation de trous noirs dotés d'un champ électromagnétique expliquerait leur origine.

Les sursauts gamma sont de brèves émissions de rayonnements gamma et X, découvertes accidentellement en 1969 par des satellites militaires américains. Ils sont souvent suivis d'une émission tardive de rayons X et de lumière visible qui peut durer de quelques jours à quelques mois. Étant donné leur fréquence et leur répartition uniforme sur le fond du ciel, les astronomes pensent qu'il se produit chaque jour, une explosion énorme quelque part dans tout le volume de l'Univers observable. Les sursauts gamma semblent être les images de cataclysmes assez puissants pour être détectables sur des distances de plusieurs milliards d'années-lumière, ce qui équivaut à la production, pendant quelques secondes du sursaut, d'une énergie comparable à celle qui est rayonnée par toutes les étoiles contenues dans toutes les galaxies de l'Univers observable! Nous avons élaboré un modèle selon lequel ces événements hors du commun signaleraient la naissance de trous noirs dotés d'un champ électromagnétique.

Nous avons étudié l'effondrement d'une étoile en rotation dotée d'un champ électromagnétique. Beaucoup d'étoiles, le Soleil, par exemple, sont entourées d'un champ magnétique produit par des courants électriques dus à la circulation de la matière stellaire dans leurs profondeurs. Lorsque les couches internes d'une étoile massive se contractent, les lignes de force de son champ électromagnétique suivent le mouvement de la surface à laquelle elles sont ancrées et se resserrent rapidement. En d'autres termes, le champ est «concentré». Lorsque l'étoile forme un trou noir, cette amplification est énorme, parce que le diamètre de la sphère de matière à laquelle le champ est attaché passe de quelques dizaines de millions de kilomètres à environ un kilomètre.

Or, en 1935, Werner Heisenberg montra qu'au-delà d'une valeur critique, un champ électromagnétique doit provoquer un phénomène qu'il nomma polarisation du vide. En effet, en mécanique quantique, on montre que le vide est, en permanence, rempli de paires de particules et d'antiparticules qui surgissent spontanément, avant de s'annihiler mutuellement. On qualifie ces particules de «virtuelles» parce qu'elles disparaissent presque immédiatement après leur formation. Cependant, en présence d'un champ électromagnétique intense, les paires d'électrons (de charge négative) et d'antiélectrons (de charge positive) sont durablement séparées par l'action du champ et ne s'annihilent pas immédiatement. De virtuelles, ces particules deviennent réelles. Ainsi, lorsqu'un trou noir doté d'un champ électromagnétique se forme, il est entouré d'une sphère, nommée dyadosphère (du grec *duas* qui signifie paire) à l'intérieur de laquelle, le champ électromagnétique est capable de créer une énorme quantité d'électrons et d'antiélectrons en quelques cent milliardièmes de nanoseconde.

En 1975, l'un d'entre nous (Remo Ruffini) et Thibault Damour, de l'Observatoire de Paris, ont démontré que ce phénomène est le plus efficace pour extraire de l'énergie d'un trou noir. Contrairement aux étoiles, dont seuls quelques pourcent de la masse sont convertis en énergie au cours de leurs dix milliards d'années d'existence, la dyadosphère entou-

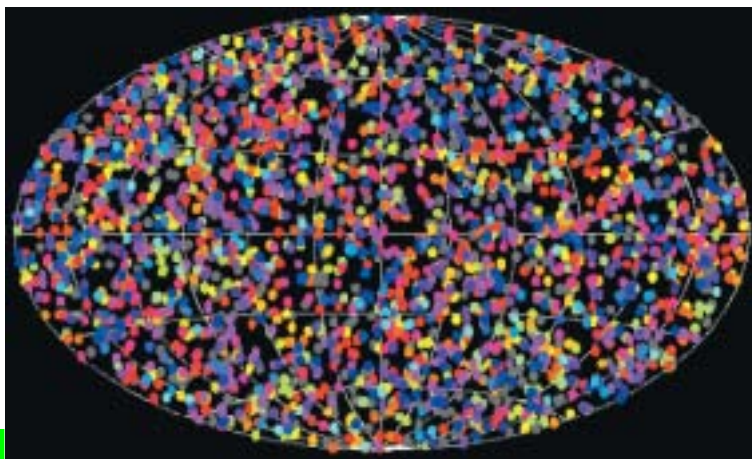
rant un trou noir doté d'un champ électromagnétique peut, en quelques secondes, convertir en énergie jusqu'à 50 pourcent de la masse de l'étoile qui lui a donné naissance.

La dyadosphère entre ensuite en expansion sous la forme d'une coquille de plasma formé d'électrons, d'antiélectrons et de photons gamma (produits par les collisions entre particules et antiparticules) que nous avons nommés impulsions électromagnétiques en paires. Expulsée à des vitesses proches de celle de la lumière, cette impulsion entre ensuite en collision avec les couches supérieures de l'étoile, essentiellement constituées d'hydrogène, qui n'avaient pas été absorbées par le trou noir. Ces restes enrichissent l'impulsion, en électrons, protons et neutrons, jusqu'à ce qu'elle arrive dans une région où le plasma est assez dilué pour être transparent. Les photons gamma sont libérés brusquement et produisent le sursaut proprement dit. Par la suite, les protons, et autres particules lourdes accélérées par l'impulsion, viennent frapper le gaz raréfié du milieu interstellaire et provoquent un gigantesque feu d'artifice en rayons X, en lumière visible et en rayonnement radio responsable de l'émission tardive.

Nous avons comparé ce modèle avec les données fournies par trois observatoires spatiaux d'astronomie (BATSE, XMM-Newton, et Rossi-XTE tous trois américains) pour un sursaut gamma intervenu le 12 décembre 1999. Les calculs s'accordent parfaitement aux observations.

Contrairement à la plupart des modèles proposés jusqu'ici pour expliquer les sursauts gamma, le nôtre ne dépend que de deux paramètres libres : l'énergie de la dyadosphère et la densité de matière dans les couches de l'étoile qui restent à l'extérieur du trou noir. Tous les autres (température, vitesse, densité du plasma, etc.) découlent directement de la théorie et sont parfaitement connus. Il est remarquable que ces prévisions soient valables aussi bien pour le sursaut proprement dit (qui se déroule sur des durées de l'ordre de la milliseconde et sur des distances de l'ordre du millier de kilomètres) que pour l'émission tardive (qui dure quelques mois et se produit à l'échelle de l'année-lumière). De surcroît, ce modèle explique pourquoi le satellite *Chandra* a détecté une raie spectrale du fer dans l'émission retardée de certains sursauts : elle serait produite par le fer contenu dans une étoile voisine lors de son explosion en supernova sous l'effet du choc avec le flash gamma. Il s'agirait là d'un nouveau mécanisme de formation des supernovae.

Remo RUFFINI, Carlo BIANCO, Federico FRASCHETTI, She-Sheng XUE, Université La Sapienza, Rome – Pascal CHARDONNET, IN2P3



La répartition des sursauts gamma (d'intensités différentes) sur le fond du ciel est uniforme d'après cette carte établie par le satellite BATSE. Les astronomes en déduisent qu'ils sont répartis de façon homogène dans l'Univers observable. Les sursauts sont les images de cataclysmes détectables à des milliards d'années-lumière.

Une mite exceptionnelle

La sexualité d'une mite fait tomber un mythe

Au royaume des métazoaires, la diploïdie est reine ! Expliquons-nous : la quasi-totalité des cellules d'animaux pluricellulaires disposent de deux exemplaires de chaque chromosome. Une cellule est haploïde quand elle n'a qu'un chromosome de chaque type. Ainsi, les êtres humains sont diploïdes : nos cellules ont 23 paires de chromosomes et ont chaque gène (à l'exception de certains portés par les chromosomes sexuels) en double.

Cependant, toute règle souffre des exceptions (sauf cette règle !) et les zoologistes connaissent des animaux haploïdes : les mâles sont haploïdes et les femelles diploïdes. C'est, par exemple, le cas de certains vers. Toutefois, on ne connaissait pas d'organismes pluricellulaires complètement haploïdes. Plus encore, on pensait que la diploïdie avait été très tôt sélectionnée dans le règne animal au profit de l'haploïdie. Celle-ci était condamnée à disparaître devant les atouts majeurs de la diploïdie : le brassage des gènes est favorable à l'évolution, et les « gènes de secours » remplacent parfois les gènes défaillants.

La mite *Brevipalpus phoenicis* étudiée par Andrew Weeks et ses collègues de l'Université d'Amsterdam fait exception. Les mâles, mais aussi les femelles de cette espèce, sont haploïdes. Un dogme est tombé. *Brevipalpus phoenicis* est une mite, de la famille des acarins, phytophage qui vit dans les régions tropicales et subtropicales. Par sa voracité ou les virus qu'elle transmet, elle cause d'importants dégâts dans les plantations de citrons, de café, de thé... Parmi les populations étudiées, au Brésil, les mâles sont très rares. Ces mites se reproduisent par parthénogenèse, c'est-à-dire qu'un ovule non fécondé se développe en un organisme adulte femelle. Les cellules de ces animaux ont deux chromosomes. Sont-ils différents (haploïdie) ou homologues (diploïdie) ? Pour trancher, les biologistes ont identifié la structure de ces chromosomes. Ainsi, ils ont détecté qu'une zone, nommée région de l'organisation nucléaire, n'est présente que sur un seul des deux chromosomes. Par ailleurs, une région d'ADN qui code une partie des ribosomes n'existe aussi que sur un seul des deux chromosomes. Enfin, l'étude de microsatellites (des séquences d'ADN répétées) a montré qu'aucun acarien ne porte plus d'un exemplaire d'un gène donné. Les mites *Brevipalpus phoenicis* sont bien haploïdes.

Un examen attentif des œufs a révélé qu'ils renfermaient des bactéries, dites endosymbiotiques, de la famille *Bacteroides*. Afin d'identifier leur rôle, les biologistes les ont éliminées de certains œufs à l'aide d'un antibiotique : les mites qui se sont développées étaient en majorité des mâles ! Une température d'incubation élevée est aussi une cause de formation de mâles. Ainsi, une bactérie féminise les mites de cette espèce mais on ignore encore par quels mécanismes. L'infection d'un ancêtre haplo-diploïde de *Brevipalpus phoenicis* par des bactéries féminisantes a sans doute entraîné l'apparition d'œufs complètement haploïdes qui se développent uniquement en adultes femelles. Grâce à cette nouvelle espèce, on pourra comparer *in situ* les avantages et les inconvénients de l'haploïdie et de la diploïdie.

Les mites *Brevipalpus phoenicis* des zones tropicales sont des organismes haploïdes : leurs deux chromosomes n'existent qu'en un seul exemplaire tant chez le mâle (en haut) que chez la femelle (en bas).



La biodiversité

Elle est indispensable au bon fonctionnement des écosystèmes.



La parcelle de Silwood, en Angleterre

Sauvegarder la biodiversité ! Tel est le mot d'ordre maintes fois entendu dans divers discours écologistes. Pourquoi cette biodiversité est-elle si importante ? Pour des raisons éthiques ? Esthétiques ? Pas seulement. Michel Loreau de l'École normale supérieure, et Andy Hector, du Collège impérial d'Ascot, ont montré que le nombre d'espèces végétales influe sur le bon fonctionnement des écosystèmes.

En 1996, dans le cadre de l'étude BIODDEPTH menée dans huit pays européens (l'Allemagne, la France, l'Irlande, le Royaume-Uni, la Suisse, le Portugal, la Grèce et la Suède), des écologistes ont isolé plusieurs parcelles, où ils ont éradiqué progressivement les plantes existantes, afin d'y planter des plantes à fleurs et des herbes à raison de 2 000 plants par mètre carré. Ils ont établi cinq degrés de diversité, de la monoculture à la richesse initiale du site en espèces végétales.

Malgré les différences climatiques de chaque pays, les plantes poussent d'autant mieux qu'elles sont en présence d'espèces multiples. Plusieurs caractéristiques de l'écosystème augmentent avec le nombre d'espèces : la productivité des plantes (la masse végétale à l'air libre et celle enterrée), la rétention d'éléments, tel l'azote, le recyclage de ces éléments, la résistance aux maladies et la composition du sol en invertébrés.

Toutefois, l'interprétation de ces résultats a longtemps été controversée. Certains y voyaient un effet de sélection, c'est-à-dire que plus les espèces sont nombreuses, plus la probabilité d'en avoir une dotée d'une forte productivité est grande. Selon d'autres, les espèces améliorent leur exploitation des ressources par un effet de complémentarité : par exemple, les légumineuses fixent l'azote atmosphérique qui fertilise le sol et profite ainsi à d'autres espèces.

M. Loreau et A. Hector ont repris les résultats de l'étude BIODDEPTH et les ont analysés à l'aune de méthodes originales inspirées des techniques statistiques utilisées en génétique afin de distinguer l'influence de chaque effet. Ils en concluent que la moyenne de l'effet de sélection sur les différents sites est nulle, et que la relation entre le nombre d'espèces et l'augmentation de productivité sur l'ensemble des parcelles ne dépend que de l'effet de complémentarité.

Les écologistes ont ainsi montré expérimentalement ce qu'ils ont longtemps été affirmé sans preuve, à savoir que le fonctionnement d'un écosystème dépend de la variété des espèces présentes dans cet écosystème. Toutefois, des questions demeurent. Par exemple, combien d'espèces participent à cet effet de complémentarité ? Suffit-il de quelques espèces seulement, ou un grand nombre doivent-elles interagir pour assurer le niveau de productivité observé dans les écosystèmes naturels ? Quoiqu'il en soit, les autorités compétentes dans la protection de l'environnement et la sauvegarde de la biodiversité disposent désormais d'un argument solide pour se faire entendre.

La panique modélisée

In modèle reproduit l'écoulement des foules paniquées en assimilant les individus à des particules mues par des «forces sociales»

A Johannesburg, en avril 2001, plus de 120 000 billets ont été vendus pour le match opposant les *Kaizer Chiefs* aux *Orlando Pirates*. L'équipe locale de football américain devait affronter sa concurrente venue de Floride dans le stade, prévu pour 68 000 personnes. Lorsque des spectateurs munis de billets, mais bloqués dehors ont commencé à pousser pour rentrer, la panique s'est répandue dans le stade. Plus de 120 personnes sont mortes piétinées, des centaines ont été blessées.

La question de l'écoulement des foules préoccupe tous les organisateurs de grands rassemblements, et les mathématiciens tentent de modéliser ces mouvements. Tous les modèles du passé sont fondés sur l'idée qu'une foule fuit la menace, comme une masse fluide et homogène s'écoule sous la pression. Toutefois, les molécules qui composent un fluide ignorent la peur et les comportements sociaux... Pour tenir compte de ces caractéristiques humaines, Dirk Helbing et ses collègues de l'Université de Dresde ont reformulé le modèle classique en représentant par des forces les principaux instincts humains en groupe. Il reproduit ainsi beaucoup mieux la fuite d'une foule paniquée.

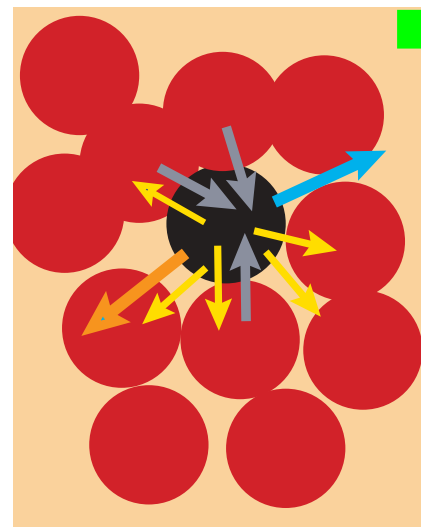
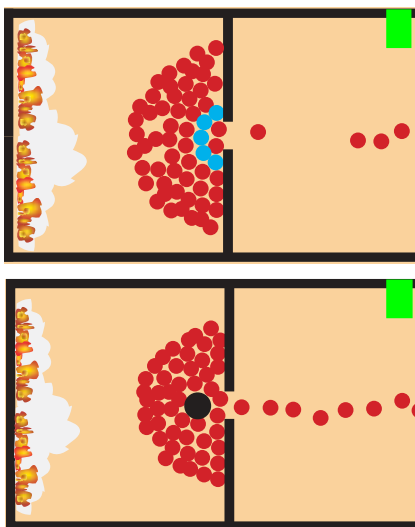
D. Helbing a mis au point une «mécanique pour particules bipèdes». Il suppose que l'accélération de chacun des piétons qui compose la foule est égale, à chaque instant, à la somme des forces qu'il subit. Elles sont de trois types : une «force motrice» qui traduit le fait que toute personne dirige sa vitesse en fonction d'objectifs propres ou d'instincts ; en cas de panique, cette force contient une composante grégaire dirigée vers ce que le groupe a reconnu comme le salut (par exemple, la sortie) ; une «force psycho-sociale» répulsive à courte distance qui

pousse chaque membre d'une foule à maintenir une distance minimale entre lui et ses voisins ; finalement deux «forces de contact», une force «corporelle», que chaque piéton exerce sur ses voisins pour contrecarrer la pression qu'il subit et l'autre de «friction» qu'il exerce lorsqu'il se faufile. À ces forces entre individus s'ajoutent les forces entre les piétons et les murs traitées de façon similaire.

Dans une première simulation, D. Helbing a étudié comment une foule de 200 personnes évacue une pièce carrée de 15 mètres de côté par une ouverture de un mètre de largeur. Tant que tout est normal, seule une légère concentration de gens se forme en face de la porte. Chaque individu progresse vers la sortie de un mètre par seconde environ, vitesse qui augmente jusqu'à 1,5 mètre par seconde si une raison de se presser apparaît, et beaucoup plus encore si la panique se répand. En présence de danger, les personnes enfermées dans la pièce commencent à pousser et à s'entrechoquer. Ces contacts augmentent les «forces de contact», notamment celles de «friction», de sorte que le mouvement général se ralentit. Celui-ci devient encore plus lent si la panique gagne

la foule. La «force grégaire» augmente alors, car chacun veut fuir vite, et la concentration d'individus en face de la porte provoque la formation d'«arcs-boutants» : des structures d'où les individus ne s'extraient qu'avec difficulté. Pour rechercher une solution, l'équipe de Dresde a montré, dans une deuxième simulation, que dans une telle situation, la vitesse de sortie augmente notablement quand une colonne, placée en face de la porte, empêche la formation des arcs-boutants.

Une troisième simulation met en évidence le rôle contrasté de la «force grégaire». Elle reproduit l'évacuation d'une pièce enfumée munie de deux sorties. Comme les individus enfermés ont des difficultés à localiser la sortie, ils commencent par la rechercher. Puis, dès qu'une première personne la trouve, ils s'engouffrent à sa suite. Seuls, quelques individus profitent de la découverte plus tardive d'une autre personne qui a trouvé la seconde sortie : leur fuite est plus rapide. Le modèle confirme que l'instinct grégaire est à double tranchant. S'il permet à une foule de profiter du salut découvert par un seul, il peut aussi la précipiter dans l'abîme.



Lorsque sous la pression d'une menace (le feu), une foule cherche à quitter une pièce, elle s'agglutine près de la sortie, et des «arcs-boutants» (a, en bleu) se forment entre les individus, ce qui ralentit l'évacuation. Une colonne (b, en noir) judicieusement placée en face de la porte empêcherait la formation des arcs-boutants qui obstruent la sortie. Dans un nouveau modèle, on représente les interactions entre individus par quatre forces : une force motrice (c, en bleu), une force psycho-sociale (en jaune), et deux forces de contact : une force corporelle (en gris) et une force de friction (en orange).

Vous cherchez un article scientifique ou technique ?

[//articlesciences.inist.fr](http://articlesciences.inist.fr)

Le moteur de recherche et de commande d'articles scientifiques et techniques

Un service de l'Institut de l'Information Scientifique et Technique du CNRS

www.inist.fr

Une sexualité de survie

Menacés d'extinction, les cyprès de Duprez ont recours à une sexualité particulière.

Au Sud-Est de l'Algérie, près de la frontière lybienne, dans les montagnes du Tassili des Ajjer, les 231 cyprès de Duprez (*Cupressus dupreziana*) sont les derniers survivants de cette espèce. Dans une telle population restreinte, le danger de la consanguinité et l'accumulation de mutations délétères qui en résulte. Le cyprès de Duprez est-il pour autant condamné? Non, l'équipe de Christian Pichot, de l'Institut national de la recherche agronomique, à Avignon, a mis en évidence chez ces arbres une sexualité unique, où le grain de pollen donne, seul, naissance à un nouvel individu.

Plusieurs indices ont mis les botanistes sur la voie. D'abord, les grains de pollen du cyprès de Duprez sont plus gros que ceux des autres espèces de cyprès. Par ailleurs, les fruits de croisements entre ce cyprès et celui de Provence ressemblaient beaucoup au cyprès de Duprez lorsque celui-ci est le père. Comment expliquer ces observations?

C. Pichot et ses collègues ont analysé les gros grains de pollen du cyprès de Duprez et les individus issus de croisements. Le grain de pollen est diploïde, c'est-à-dire qu'il contient deux lots de chromosomes paternels, alors que ceux d'autres espèces sont haploïdes (ils n'ont qu'un seul lot de chromosomes). Ensuite, les cyprès issus des croisements du pollen de *C. dupreziana* et de l'ovule d'une autre espèce sont génétiquement identiques à l'arbre père. Ainsi, le pollen se pose sur l'ovaire, émet un tube pollinique qui introduit le noyau dans l'ovule, mais la fusion des deux noyaux n'a pas lieu : celui de l'ovaire dégénère alors que celui, déjà diploïde, du grain de pollen prend en charge le développement du nouvel arbre. L'ovule n'est alors qu'une réserve nourricière.

Ce mode de reproduction est nommé apomixie. Il est connu chez les angiospermes (les plantes à fleurs), chez qui le développement d'un ovule seul en un nouvel individu est fréquent. Cependant, l'apomixie du cyprès de Duprez est inédite à deux titres : elle concerne les gamètes mâles, le pollen, et apparaît chez un gymnosperme, une plante à brins nus sans fleurs.

L'apomixie apporte un avantage aux cyprès qui, grâce à elle, évitent la consanguinité ; cependant, elle freine leur évolution, car elle supprime le brassage génétique. L'adoption d'un tel mode de reproduction rend délicate la sauvegarde de l'espèce menacée à long terme : les généticiens de l'INRA espèrent analyser chacun des 231 survivants dans l'espoir de découvrir un exemplaire encore capable de fécondation. Enfin, ils étudient différents croisements interspécifiques qui sauvegarderaient le patrimoine génétique des cyprès de Duprez, et, pourquoi pas, ramèneraient dans le chemin de la fécondation.



Dans le Sud-Est de l'Algérie, les cyprès de Duprez résistent à l'extinction grâce à une sexualité originale : les grains de pollen, à eux seuls, donnent naissance à de nouveaux cyprès.

L'homme à face plate

Un nouveau visage dans notre arbre généalogique.

Après avoir trouvé *Errorin tungensis*, un hominidé de six millions d'années.

Les paléontologues ont annoncé, quelques mois plus tard, la découverte d'un autre de nos ancêtres au Kenya. Le crâne mis au jour est presque complet :

Kenyanthropus platyops l'homme à face plate.

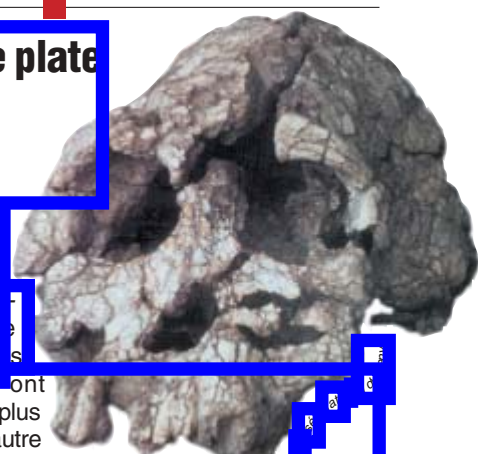
Il daterait de 3,5 millions d'années environ. Toutefois, ses caractères à la fois primitifs et évolués laissent les paléontologues perplexes quant à la place qu'il occupe dans notre arbre généalogique.

Le fossile, mis au jour par l'équipe de Meave Leakey, paléontologue aux Musées nationaux du Kenya, se trouvait à Lomekwi sur la rive Ouest du lac Turkana, dans le Nord du Kenya. Ce site a livré plus de 30 fragments de crânes et de dents, dont deux ont été attribués au nouvel hominidé (les autres n'ont pas encore été classés). Tous ces fragments ont été découverts entre deux couches de sédiments datées de 3,2 et 3,5 millions d'années. D'après les restes de mammifères qui les accompagnaient, la région était recouverte, à cette époque, de forêts et de prairies, comme les sites de Tanzanie et d'Éthiopie où l'on a retrouvé des restes d'australopithèques.

Le nouveau fossile présente une mosaïque de caractères qui ne sont pas nouveaux en soi, mais qu'aucun autre fossile ne rassemble. Par le petit volume de son crâne, le nouvel hominidé ressemble au chimpanzé, mais par ses dents à émail épais (qui révèlent la dureté des aliments dont il se nourrissait) il se rapproche des australopithèques de l'espèce *afarensis* vieille de 3,2 millions d'années (celle de Lucy) et de l'espèce *anamensis*, encore plus ancienne. Enfin, par sa face plate, il ressemble à *Homo rudolfensis*, découvert sur la rive opposée du lac Turkana, mais beaucoup plus jeune, puisqu'il daterait de 2 à 2,5 millions d'années. Ainsi, ne pouvant attribuer le nouveau fossile à aucun de nos ancêtres déjà connus, les paléontologues ont classé dans un nouveau genre, qu'ils ont nommé *Kenyanthropus platyops*, c'est-à-dire l'homme à face plate.

Ses caractères propres le distinguent des australopithèques et des paranthropes (un genre d'hominidés aussi appelés les australopithèques «robustes», mais plus jeunes que les australopithèques, puisqu'ils prospéraient il y a 2,5 à 1,4 millions d'années). Qu'en est-il de sa ressemblance avec *Homo rudolfensis*, ou plus précisément avec le fossile dit de «l'homme 1 470», daté de 1,8 millions d'années, à partir duquel on a défini ce genre? Certains paléontologues pensent que *Homo rudolfensis* est plus proche des australopithèques que des *Homo*. Au lieu de créer un nouveau genre, Meave Leakey aurait pu considérer que les paranthropes et les *Homo* primitifs étaient en fait des australopithèques, mais ce dernier genre serait alors devenu un «fourre-tout». Elle suggère plutôt de classer «l'homme 1 470» dans le groupe des kenyanthropes (celui de *Kenyanthropus platyops*). Cette proposition controversée ravive des désaccords déjà anciens à propos de ce fossile, retrouvé en petits morceaux, de surcroît écrasés par le poids des sédiments.

Comment expliquer les ressemblances de *Kenyanthropus platyops* avec des hominidés si différents les uns des autres et si éloignés dans le temps? Les indices dont nous disposons (les



os et les dents) pourraient davantage tenir de l'adaptation à l'environnement que de l'évolution. Depuis une vingtaine d'années, les découvertes de fossiles s'accumulent et la famille de l'homme s'agrandit. Cependant, ses origines ne s'éclaircissent pas pour autant : l'ancienne théorie, selon laquelle les hominidés s'étaient succédé selon un processus d'évolution linéaire

n'a plus cours. On avait d'abord pensé que les genres d'hominidés plus proches de l'homme que du chimpanzé n'étaient qu'un nombre de trois, à savoir les australopithecus, les paranthropes et les *Homo*. Puis, la découverte de trois nouveaux australopithecus, au Kenya, en Éthiopie et au Tchad a montré que ce genre était plus répandu qu'on ne le supposait. Depuis

deux autres genres ont été créés : *Ardipithecus*, d'abord daté de 4,4 millions d'années, mais qu'un fossile récemment trouvé fait remonter à 5,8 millions d'années, et *Orrorin*, de 6 millions d'années. La découverte de *Kenyanthropus platyops*, qui ajoute à la confusion, laisse supposer que les incertitudes qui planent sur nos origines ne sont pas près de disparaître.

Cerveau et pesanteur

Pour prévoir la trajectoire d'un objet, notre cerveau tient compte de l'accélération due à la pesanteur.

Le petit d'homme acquiert très rapidement la notion de la pesanteur : l'enfant de un an sait que s'il lâche un objet, celui-ci tombe. Jusqu'à que

la question se pose notamment lorsque nous attrapons un objet qui nous est lancé. Sachant que les influx nerveux mettent un certain temps pour aller du cerveau aux muscles et que ces derniers ne répondent pas immédiatement, le cerveau doit donner les ordres pour la saisie de l'objet (lever l'avant-bras, contracter les muscles du biceps, puis tendre les doigts) de quelques dizaines à quelques centaines de millisecondes avant que l'objet n'arrive au niveau de la main. En d'autres termes, il doit anticiper la trajectoire afin que la main se trouve au bon endroit, au bon moment. Comment prévoit-il la trajectoire ? Une équipe du Laboratoire européen de neurosciences de l'action, dont Joe McIntyre et Alain Berthoz du Collège de France, à Paris, et Francesco Lacquaniti et Myrka Zago, à Rome, ont comparé des saisies de balles réalisées sur Terre et en impesanteur au cours d'un vol de la navette américaine. Ils ont montré que le cerveau prévoit une trajectoire en tenant compte de l'accélération de la pesanteur.

On peut envisager *a priori* trois hypothèses sur la façon dont le cerveau opère pour prévoir une trajectoire. Selon la première, le cerveau évalue la vitesse de l'objet et ne se fonde que sur cette information. Dans la deuxième, il estime non seule-

ment la vitesse, mais aussi l'accélération de l'objet. Dans la troisième, il n'apprécie que la vitesse, mais sait que l'objet est soumis à la pesanteur et que cela modifie sa trajectoire. Le système visuel évalue mal les accélérations, de sorte que l'ordre éliminer la deuxième hypothèse.

Pour trouver laquelle des deux autres hypothèses est la bonne, J. McIntyre et ses collègues ont filmé le mouvement du bras de spationautes qui devaient attraper une balle lancée vers eux, par le haut, à différentes vitesses. Afin que les sujets ne s'accoutument pas à la saisie de la balle, la vitesse des lancers variait de façon aléatoire. La même expérience a été réalisée également sur Terre.

Sur Terre, l'avant-bras se lève environ 200 millisecondes avant l'impact. Dans l'espace, il se lève plus tôt (500 millisecondes avant l'impact pour une vitesse de lancer égale à 1,7 mètre par seconde) et arrête sa course, voire l'inverse, quand le sujet s'aperçoit qu'il n'attrapera pas la balle. L'avant-bras a manifestement reçu trop tôt l'ordre de bouger. Ce résultat écarte d'emblée la première hypothèse : dans le cas où le cerveau ne se fonderait que sur la vitesse de la balle pour prévoir la trajectoire, l'homme capturerait la balle au bon moment dans l'espace, où la vitesse est constante.

Reste la troisième hypothèse, celle de la prédiction qui tient compte de la pesanteur. Dans le cas où le cerveau tiendrait compte de la pesanteur, il s'attend à ce que la balle soit accélérée et, par conséquent, plus basse qu'elle ne l'est en réalité dans l'espace : il sous-estime la durée avant impact et déclenche des mouvements trop tôt. Cette sous-estimation est d'autant plus importante que la vitesse de lancer est faible. C'est bien ce que J. McIntyre et ses collègues ont observé : le retard pris dans

l'espace passe de 100 millisecondes pour une vitesse de lancer de 2,7 mètres par seconde à 700 millisecondes pour 0,7 mètre par seconde. On explique quantitativement ces variations en considérant que le cerveau prévoyait que l'objet allait subir une accélération égale à celle de la pesanteur terrestre et qu'il s'est décidé à donner l'ordre à l'avant-bras de bouger environ 300 millisecondes avant ce qu'il pensait être l'impact. J. McIntyre et ses collègues concluent que le cerveau possède un modèle interne de trajectoire qui a assimilé la valeur de la pesanteur terrestre.

A mesure des expériences pratiquées dans l'espace à plusieurs jours d'intervalle, les spationautes ont montré un début d'adaptation à l'absence de gravité. Qu'est-ce qui a provoqué cette adaptation ? Des indices visuels ? Des expériences où les indices visuels feront croire aux sujets qu'ils sont tête en bas devraient répondre en partie à ces questions.



Un astronaute doit saisir une balle lancée vers lui en impesanteur à des vitesses initiales différentes.

Vous cherchez des services personnalisés d'information scientifique ?

//connectsciences.inist.fr

Le portail CNRS d'information scientifique et technique

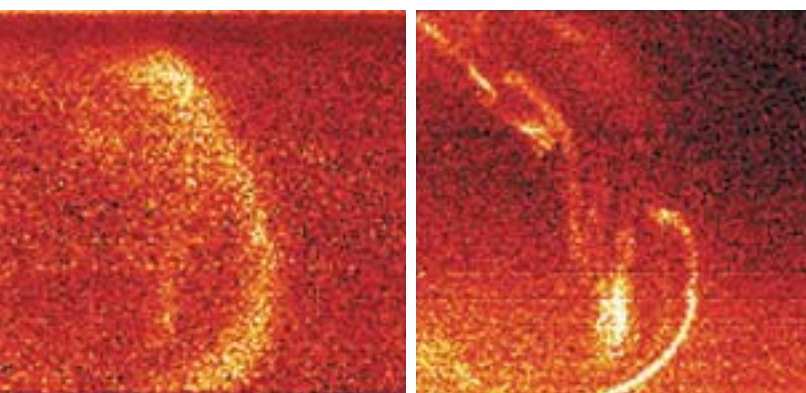
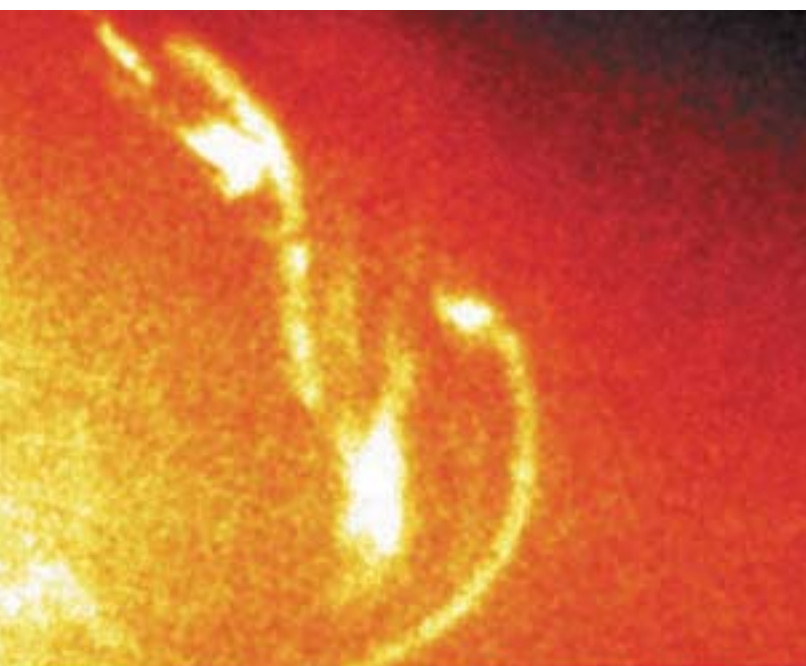
Un service de l'Institut de l'Information Scientifique et Technique du CNRS
www.inist.fr

Les aurores à protons

Pour la première fois, des aurores boréales à protons ont été observées depuis l'espace

Les aurores boréales sont d'une beauté féerique. Pourtant, une partie du spectacle, invisible à l'œil nu, nous échappe : car certaines des longueurs d'onde émises se situent dans l'ultraviolet, comme celles émises par les aurores dites à protons. Grâce à une caméra sensible au rayonnement ultraviolet, embarquée à bord du satellite américain IMAGE, en orbite depuis mars 2000, les géophysiciens du Laboratoire de physique atmosphérique et planétaire de l'Université de Liège et leurs collègues américains ont, pour la première fois, observé ces aurores à protons sur l'ensemble de l'hémisphère Nord.

Les aurores boréales classiques sont provoquées par le vent solaire qui «se heurte» à la magnétosphère terrestre.



Cette aurore boréale a été observée dans l'ultraviolet à l'aide de plusieurs caméras du satellite IMAGE : l'une des images correspond à l'azote excité principalement par les électrons auroraux (*en haut*), une autre à l'oxygène, également dû aux électrons (*en bas à droite*) ; une image de l'aurore à protons (*en bas à gauche*) a été obtenue par une caméra sensible à l'émission de l'hydrogène. La différence de structure est flagrante : l'émission aurorale est nettement plus diffuse pour les protons, et certaines régions excitées par les électrons ne le sont pas par les protons.

Le vent solaire est un plasma composé d'électrons, de protons, de noyaux d'hélium et d'ions, issus des couches externes du Soleil. Les électrons y sont presque aussi nombreux que les protons. Arrivé en périphérie de la Terre, le vent solaire rencontre la magnétosphère, qui s'étend à plus de 500 rayons terrestres, du côté opposé au Soleil, mais seulement à dix rayons terrestres du côté exposé au Soleil. Certaines particules chargées s'accumulent dans un réservoir, le «feuillet de plasma», situé dans la queue de la magnétosphère, du côté opposé au Soleil. D'autres, du côté exposé au Soleil, suivent les lignes de force du champ géomagnétique qui convergent vers les régions polaires. Quand elles abordent les couches les plus élevées de l'atmosphère terrestre, elles entrent en collision avec les molécules et les atomes de l'ionosphère (de l'azote et de l'oxygène). Elles les excitent, et quand les atomes se dés excitent, ils émettent des photons : c'est la lumière des aurores boréales.

A beaucoup d'égards, les aurores à protons sont semblables aux aurores boréales à électrons. Les protons et les électrons ont une énergie similaire, de quelques milliers d'électronvolts. Ils excitent les mêmes atomes et provoquent les mêmes émissions, à l'exception d'une excitation de l'atome d'hydrogène. Dans ce cas, le mécanisme est un peu différent. Les protons du vent solaire sont des atomes d'hydrogène qui ont perdu un électron. Quand un tel ion positif entre en collision avec les atomes des hautes couches de l'atmosphère, il capture un électron, et devient neutre pendant une fraction de seconde. L'atome d'hydrogène ainsi créé est le plus souvent dans un état excité, et il se dés excite en émettant un photon dans l'ultraviolet. Lors d'une nouvelle collision, il perd à nouveau son électron, puis en regagne un autre lors de la collision suivante, le perd de nouveau, et ainsi de suite. Un proton incident perd et retrouve successivement sa charge plus d'une centaine de fois, et émet ainsi des dizaines de photons ultraviolets.

Grâce au satellite IMAGE, les physiciens étudient la morphologie globale des aurores à protons de l'hémisphère Nord. Généralement, les électrons et les protons proviennent des mêmes régions de la magnétosphère et ils sont canalisés par les lignes de force : les aurores à électrons et à protons se superposent à environ 80 pour cent. Toutefois, chaque jour, en fin d'après-midi, l'aurore à protons se décale vers le Sud, et les deux aurores sont séparées. En effet, comme les particules ont une charge opposée, et une masse très différente (les protons sont environ 1 850 fois plus lourds que les électrons), leur comportement dans le champ électromagnétique terrestre diffère.

Les physiciens étudient les mécanismes qui précèdent les aurores : comment les particules pénètrent dans le champ terrestre, comment elles sont «triées» par les champs, et comment elles acquièrent leur énergie dans la magnétosphère, souvent très supérieure à celle qui règne dans le vent solaire. L'énergie ainsi déversée dans l'atmosphère peut dépasser 10 milliards de watts lors des fortes tempêtes magnétiques, qui suivent certaines éruptions solaires fréquentes en ce moment où le Soleil traverse une période d'activité maximale. En moyenne, les protons seraient responsables d'environ 15 pour cent de cette énergie, mais cette fraction est variable. La contribution énergétique des protons dans les aurores était considérée comme faible. Or, les astrophysiciens liégeois ont découvert que les protons sont dominants dans certaines zones aurorales, car ils transportent alors l'essentiel de l'énergie injectée dans notre atmosphère. Ils jouent un rôle primordial dans l'ionisation de la haute atmosphère et dans la structure des courants qui y circulent. Dans les aurores boréales aussi, l'essentiel est parfois invisible pour les yeux.

Un trou dans la banquise

La banquise polaire australe se trouve périodiquement sous l'effet d'un courant d'eau circulaire dû à un relief sous-marin.

Chaque année, le froid de l'hiver austral installe autour du continent antarctique une ceinture de glace qui s'étend sur plusieurs centaines de kilomètres. Toutefois, dans la mer de Weddel qui borde l'Antarctique et l'océan Atlantique, une zone de l'océan ne gèle pas systématiquement, et dessine un trou dans la banquise. Selon David Michael Holland, du Centre des sciences de l'atmosphère et de l'océan de l'Université de New York, de la glace se forme dans cette région, mais les plaques ne s'agglomèrent pas en banquise, car un courant marin circulaire les éloigne les uns des autres. Ce courant naît de la rencontre des masses d'eau qui font le tour de l'Antarctique et du mont Maud, une montagne sous-marine.

Depuis une trentaine d'années que la région antarctique est observée par satellite, la formation d'une « polynye » (mot d'origine russe qui désigne une zone dépourvue de glace dans une mer gelée) est un événement rare. C'est en 1974 que l'on a constaté la polynye la plus grande avec une absence de glace sur une superficie proche de celle de la Grande-Bretagne.

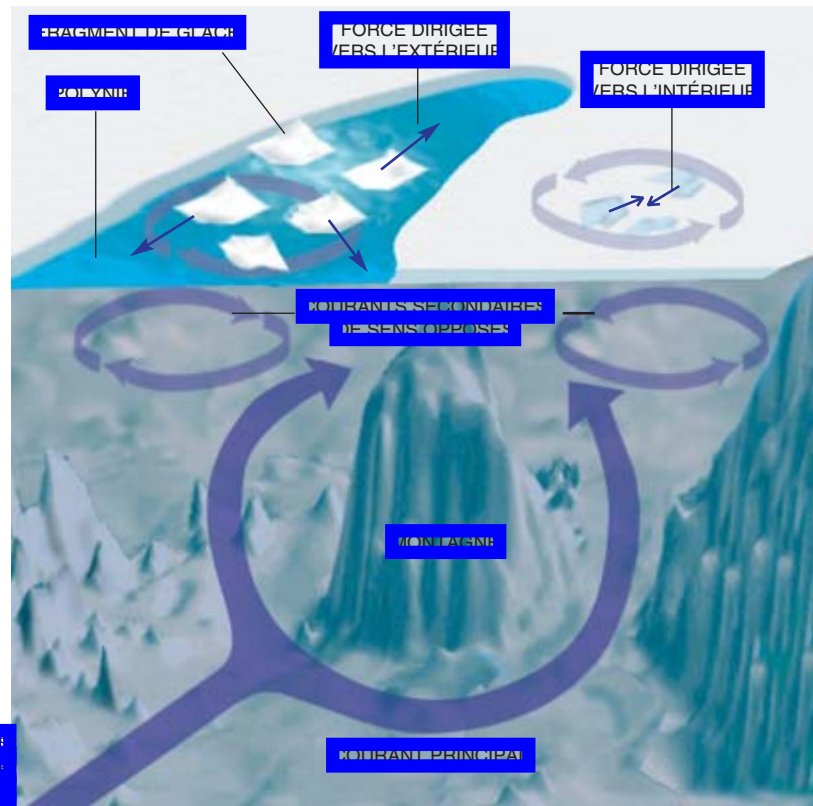
Pour expliquer la polynye de la mer de Weddell, on proposa que l'atmosphère ou des courants marins réchauffent la glace. D.M. Holland a montré que la calotte ne se forme pas non en raison d'un apport de chaleur, mais à cause d'une force qui éloigne les blocs de glace les uns des autres. Il a simulé l'écoulement du courant circumpolaire Antarctique (un courant marin qui tourne autour du continent Antarctique) en interaction avec l'atmosphère dans la région du mont Maud, une montagne sous-marine de 3 000 mètres de hauteur qui s'élève d'une plaine abyssale, et dont le sommet est à 2 000 mètres sous le niveau de la mer.

Quand les courants profonds rencontrent le mont Maud, ils contournent la montagne sous-marine et entraînent avec eux les courants de surface. Deux courants de surface circulaires se forment : l'un, au Nord du mont, est cyclonique (il tourne dans le sens de rotation de la Terre) et l'autre, au Sud, est anticyclonique. Les cristaux de glace qui naissent dans l'eau sont entraînés par ces courants, mais ne suivent pas un mouvement purement circulaire, car la force de Coriolis les repousse, selon le sens de rotation du courant, soit vers l'extérieur du courant, soit vers son centre.

La force de Coriolis résulte du mouvement de rotation de la Terre et agit sur tout objet à la surface du globe. Le sens de la force de Coriolis dépend de l'hémisphère où l'on se trouve : un objet qui se déplace dans l'hémisphère Nord est poussé vers sa droite et dans l'hémisphère Sud vers sa gauche. Revenons aux icebergs qui se forment en mer de Weddell.

Pour ceux qui se forment dans le courant cyclonique, la force de Coriolis les entraîne vers la périphérie du courant, tandis que dans le courant anticyclonique, elle agit en sens inverse et repousse les blocs de glace vers le centre du courant. Il en résulte que dans le courant cyclonique, la surface gèle, mais les plaques de glace sont poussées vers l'extérieur du courant, et laissent une zone d'eau libre de glace. Au contraire, dans la zone où les blocs sont repoussés vers le centre du courant, toute la zone finit par geler.

Les simulations réalisées par D. M. Holland reproduisent avec précision la formation des polynyes. Toutefois, une question reste ouverte : pourquoi une polynye aussi large que celle observée en 1974 n'est-elle pas réapparue depuis ? Selon D. M. Holland, la réponse tiendrait aux autres facteurs qui influent sur la formation de la polynye et qui, eux, évoluent d'une année sur l'autre : la température de l'océan, la force du courant circumpolaire ou encore celle des vents.



Un courant océanique est perturbé par la présence d'une haute montagne sous-marine. Il se sépare en deux courants secondaires de sens opposés. En surface, ces courants perturbent la formation de la couche de glace. L'un d'eux (à gauche) tourne dans le sens de rotation de la Terre et, en raison de la force de Coriolis, les blocs de glace qui s'y forment sont repoussés vers l'extérieur. Au contraire, l'autre courant (à droite) tourne dans le sens inverse de celui de rotation de la Terre et la force de Coriolis repousse vers le centre les blocs de glace qui fusionnent. Ce phénomène explique la présence, certaines années, d'un trou, ou polynye, dans la banquise antarctique.

Vous cherchez de l'Information Scientifique et Technique ?

[//www.inist.fr](http://www.inist.fr)

L'Institut de l'Information Scientifique et Technique du CNRS

Nos ancêtres, les cannibales

TIM WHITE

Les preuves de cannibalisme sont rares, mais il est désormais avéré que cette pratique est ancrée dans notre histoire.



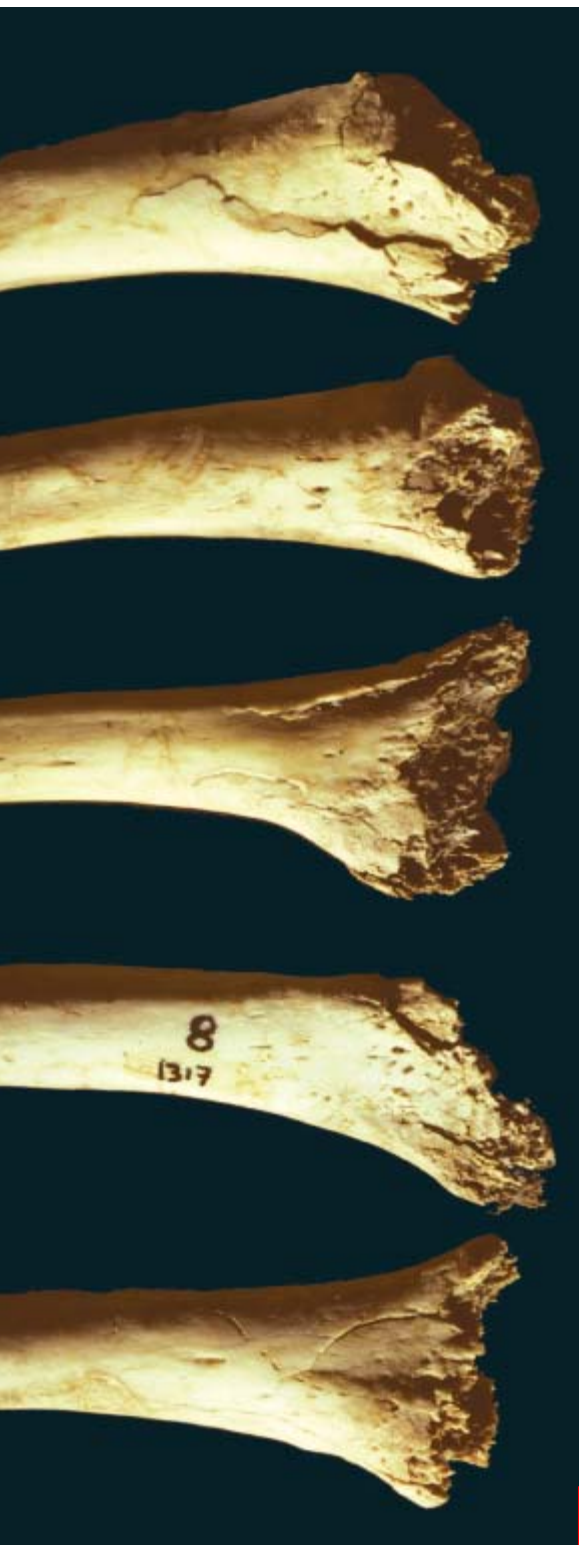
1. CE CRANE NEANDERTALIEN, mis au jour dans la grotte de Krapina en Croatie, présente des traces de cannibalisme : la calotte crânienne a été brisée pour en extraire le cerveau et le consommer. Sur ce site, des centaines d'autres ossements portent aussi des marques de cannibalisme (traces de découpe, os brisés...)



Le cannibalisme choque, effraie et fascine, qu'il s'agisse de récits de pionniers affamés, de rescapés d'accidents d'avions qui se nourrissent de leurs congénères décédés, ou de rites pratiqués en Papouasie Nouvelle-Guinée. Il représente l'un des derniers tabous et une pratique que l'on associe à d'autres cultures, à d'autres lieux et à d'autres époques. Jusqu'à présent, les éléments rassemblés depuis quelques siècles étaient trop confus et trop incomplets pour que nous puissions affirmer que cette pratique n'avait jamais existé, ou, au contraire, qu'elle avait été en vigueur à tel endroit et à telle période. De récentes découvertes ont apporté un éclairage nouveau sur le cannibalisme. Bien avant la découverte des métaux, avant la construction des pyramides d'Égypte, avant l'apparition de l'agriculture ou de l'art pariétal au Paléolithique supérieur, le cannibalisme était pratiqué par de nombreux peuples. Des os humains brisés et éparpillés ont été découverts par milliers sur les sites préhistoriques du

sud-Ouest des États-Unis, jusqu'aux îles du Pacifique. Grâce à des méthodes et à des outils de plus en plus perfectionnés, les études ostéologiques et archéologiques de ces sites ont apporté des preuves de l'existence d'un cannibalisme préhistorique.

Le cannibalisme humain intrigue depuis longtemps les anthropologues qui ont établi une classification du phénomène, par exemple en fonction du sujet consommé. Ainsi, l'endocannibalisme est la consommation d'individus à l'intérieur d'un groupe, l'exocannibalisme celle d'individus extérieurs au groupe, et l'autocannibalisme va des ongles rongés à l'autoconsommation sous la torture. Les anthropologues ont aussi établi des classifications en fonction des motivations possibles ou avérées. Par exemple, le cannibalisme parmi les survivants d'un accident est motivé par la faim. Ce fut le cas des naufragés du radeau de La Méduse ou des joueurs de rugby perdus dans les Andes, sans nourriture, en 1972, après un accident d'avion. Le cannibalisme



2. L'ECRASEMENT est un des signes de la consommation d'os humains. Les os consommés par les cannibales présentent des traces variées. Lorsqu'elles sont identiques à celles relevées sur les ossements d'animaux provenant des mêmes sites, les archéologues supposent que les restes humains ont été préparés de la même manière et pour les mêmes raisons : pour être mangés. Ces cinq métatarsiens (une partie du squelette du pied) proviennent du Canyon Mancos, dans le Colorado : l'os spongieux a été écrasé pour en extraire la graisse.

dit rituel intervient lorsque les membres d'une famille ou d'une communauté consomment leurs morts au cours de rites funéraires, afin de s'approprier leurs qualités ou d'honorer leur mémoire. Enfin, le cannibalisme pathologique est réservé aux criminels qui consomment leurs victimes ou, plus souvent, à des personnages de fiction, tel Hannibal dans *Le Silence des agneaux*.

Malgré ces distinctions, la plupart des anthropologues utilisent le terme de cannibalisme pour la consommation régulière et culturelle de chair humaine. À l'époque des grandes expéditions ethnographiques, c'est-à-dire de 400 avant notre ère avec l'historien grec Hérodote jusqu'au début du XX^e siècle, les voyageurs, les missionnaires, les militaires et les anthropologues ont exploré des contrées inconnues et ont rapporté des récits de cannibalisme alimentaire d'Amérique centrale, d'Afrique centrale ou des îles du Pacifique. Ces témoignages ont souvent été controversés, et doivent être traités avec prudence, d'autant que les anthropologues n'ont participé aux expéditions qu'à partir de la fin du XIX^e siècle. En 1937, l'anthropologue Ashley Montague déclara même que le cannibalisme n'était qu'un «mythe de voyageurs». En 1979, un autre anthropologue, William Arens de l'Université Stony Brook, à New York, nia l'existence du cannibalisme rituel. Selon lui, les témoignages de cannibalisme chez les Aztèques, les Maoris ou les Zoulous n'étaient pas fiables. On reprocha aux anthropologues de ne pas s'être limités à l'étude des populations contemporaines, et d'avoir interprété des données archéologiques à la lumière des préjugés de l'époque. En 1871, le romancier américain Mark Twain s'intéressa à ce sujet. Selon lui, quand on découvre un tas d'os d'hommes primitifs mélangé à des ossements d'animaux, sans le moindre indice qui indique que l'homme a mangé l'ours ou que l'ours a mangé l'homme, on peut y trouver des «preuves» de cannibalisme, si l'on en cherche et si l'on interprète sans scrupules les restes d'un homme décédé il y a deux millions d'années.

Durant le siècle suivant, les hominidés *Australopithecus africanus*, ainsi que *Homo erectus* et *Homo neanderthalensis* étaient considérés comme ayant été des cannibales. Selon certains, le cannibalisme aurait existé depuis près

de trois millions d'années et se serait éteint très récemment.

Toutefois, au début des années 1980, de sévères critiques s'élevèrent contre ces affirmations. L'archéologue Lewis Binford déclara que rien ne prouvait que les premiers hominidés étaient cannibales. Se fondant sur les travaux d'autres préhistoriens qui s'étaient intéressés à la composition au contexte des assemblages osseux du Paléolithique et à leurs variations, il insistait sur la nécessité de se fonder sur des expérimentations et sur des observations faites sur les civilisations contemporaines pour étudier les comportements anciens.

Diversité des rites funéraires

L'étude de cannibales contemporains serait très utile, mais, aujourd'hui, les cannibales ont presque partout disparu. Ce comportement relève désormais des sciences historiques, et l'archéologie est le principal outil de recherche sur l'existence du cannibalisme humain et sur son ampleur (le terme cannibalisme signifie «qui consomme sa propre viande» et ne s'applique donc pas seulement à l'homme. Un loup qui dévore un loup est cannibale. Pour l'homme, on devrait dire anthropophage).

Les archéologues sont confrontés à l'incroyable diversité des rituels funéraires. Les corps sont enterrés, brûlés, placés sur des échafaudages, jetés à l'eau, déposés dans des troncs d'arbres ou offerts aux charognards. Les os sont parfois déterrés, lavés, peints, enterrés en groupes ou dispersés sur des pierres. Dans certaines régions du Tibet, les futurs archéologues auront bien des difficultés à reconstituer les rites funéraires pratiqués aujourd'hui. En effet, la plupart des corps sont démembrés et donnés en pâture aux vautours et à d'autres carnivores. Les os sont ensuite ramassés, réduits en poudre, mélangés à de l'orge et à de la farine, puis de nouveau offerts aux vautours. En raison des différents rites, il est parfois difficile de distinguer le cannibalisme d'autres pratiques funéraires.

Les chercheurs ont établi des critères rigoureux pour reconnaître le cannibalisme et d'autres pratiques. Le cannibalisme est avéré lorsque les marques observées sur les restes humains sont identiques à celles que l'on retrouve sur les os d'autres ani-

naux dont la viande a été consommée. Les archéologues ont longtemps discuté de la validité de cette comparaison entre les restes humains et les restes d'animaux retrouvés sur un même site. Selon eux, les traces présentes sur un os animal, et la façon dont les os ont été déposés, indiquent que l'animal a été tué pour être consommé. Lorsque des restes humains sont découverts dans des contextes similaires, qu'ils présentent les mêmes traces, le même état de conservation et qu'ils ont été entassés de la même manière, on peut en déduire que l'on est face à des restes humains dont la chair a été consommée.

Lorsqu'un mammifère en dévore un autre, il laisse généralement des marques et des stries sur le squelette de l'animal. En effet, des tissus mous ayant une valeur nutritive, recouvrent les os des mammifères. Lorsque ces tissus sont arrachés, des marques de dents sur les os ou des fractures révèlent comment le mammifère a procédé. Lorsque les humains mangent des animaux, les marques sur les os ne proviennent pas uniquement des dents, car les carcasses sont découpées avec des outils de pierre ou de métal. On retrouve les mêmes traces sur les squelettes humains dont on a mangé la viande.

Pour établir l'existence d'une pratique cannibale, il faut retrouver les stigmates de la découpe, d'un écrasement, de fractures ou de brûlures ainsi que la présence d'os ou de parties d'os intacts. Les tissus à haute valeur nutritive, tels le cerveau ou la moelle, sont situés à l'intérieur des os et ne peuvent être extraits qu'en cassant ces os, ce qui laisse des traces. Quand ce type de traces, liées à la préparation, sont présentes sur des os humains retrouvés sur des sites archéologiques, on peut supposer que l'on est face à des pratiques cannibales. On peut valider l'hypothèse grâce à d'autres données archéologiques, notamment grâce aux restes non humains consommés sur des sites habités par le même groupe.

Ce système de comparaisons repose sur la présence de différents types de traces sur les os et sur la comparaison des contextes. Les critères d'identification de cannibalisme sont rigoureux, ainsi, la présence de marques de découpe sur les os n'est pas en elle-même une preuve. Par exemple, les os provenant d'un cimetière de soldats présenteraient des marques de baïonnette et ceux des cadavres disséqués dans les



LES ENTAILLES VISIBLES sur la partie gauche de ce fragment de tibia humain témoignent de la découpe des muscles et des tendons. Des outils étaient également utilisés pour découper la chair ou décapiter le cadavre. Les archéologues doivent toutefois faire preuve de prudence dans leurs interprétations, car toute trace de découpe n'indique pas nécessairement une pratique du cannibalisme, tant les rites funéraires sont variés.

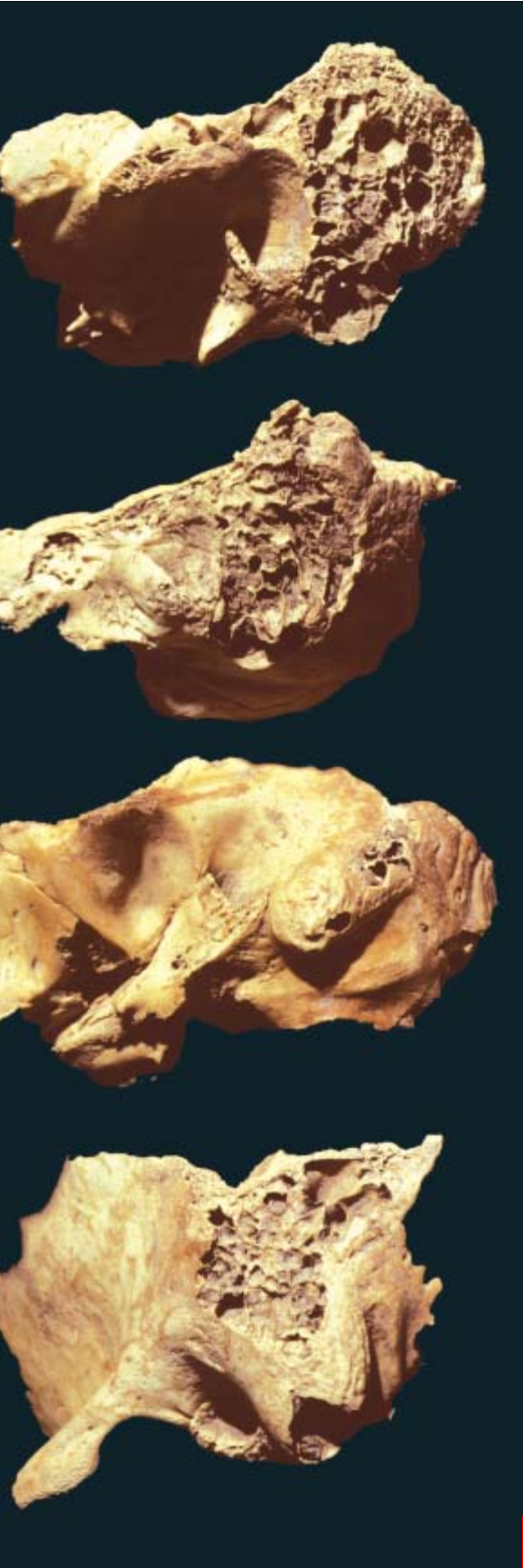
facultés de médecine des traces de découpe sans, bien sûr, qu'il s'agisse de cannibalisme. Avec des critères aussi stricts, la plupart des cas de cannibalisme ne seront pas reconnus. Des ethnographes ont rapporté un rite cannibale de Papouasie Nouvelle-Guinée : les crânes des défunts étaient soigneusement nettoyés et vidés de leur cerveau. Une fois séchés, les crânes presque intacts, étaient manipulés à maintes reprises, ce qui usait les parties saillantes. Ils étaient parfois peints et montés sur des perches pour être exposés comme des objets de culte. Les tissus mous, y compris le cerveau, étaient mangés au début des préparatifs. Cette pratique s'apparente à un cannibalisme rituel. Si ces crânes avaient été trouvés dans un contexte archéologique sans les informations dont on dispose sur ce mode de cannibalisme, ils n'auraient pas constitué des preuves de cannibalisme, selon les critères en vigueur.

Toutefois, même avec ces critères très rigoureux, nous avons découvert des indices de cannibalisme sur des sites plus anciens. Le meilleur indice de cannibalisme préhistorique provient du Sud-Ouest de l'Amérique, où les archéologues ont interprété des dizaines d'assemblages de restes humains comme des preuves incontestables de cannibalisme. D'autres preuves de pratique de cannibalisme

du Néolithique et à l'âge du Bronze ont également été trouvées en Europe, dans la grotte de Fontbregoua, dans le Var, par exemple, comme l'ont montré les travaux de Jean Courtin, du Centre de recherches archéologiques de Valenciennes, et de Paola Villa, de l'Université du Colorado. Le plus ancien site d'hominidés européens présente aussi des preuves de cannibalisme.

Les premiers Européens cannibales

Le plus important site paléanthropologique d'Europe se situe au Nord de l'Espagne, dans les contreforts de la Sierra d'Atapuerca. Ces collines regorgent de grottes, qui sont autant de sites occupés par l'homme préhistorique, mais le plus ancien découvert à ce jour en cours de fouille, est connu sous le nom de Gran Dolina. Une équipe dirigée par Eudald Carbonell, de l'Université de Tarragone en Espagne, a découvert des indices d'occupation par le nouvel ancêtre probable de l'homme, *Homo antecessor*, qui remontent à plus de 800 000 ans. Les ossements de cet hominidé ont été retrouvés dans l'une des couches sédimentaires de la grotte, mêlés à des outils de pierre et des restes de gibier (des cervidés, des bisons et des rhinocéros). On a découvert 92 fragments d'os provenant de six individus. Les



J. LES OS ÉTAIENT PARFOIS BRULÉS. Les zones endommagées sur ces quatre masoïdes (protubérance osseuse derrière l'oreille) indiquent que ces crânes humains ont été rôtis. Comme la région mastoïdienne n'est pas recouverte de beaucoup de chair, l'altération par le feu y est plus marquée que sur les autres parties du crâne. Les traces de brûlure démontrent l'existence de pratiques culinaires.

ossements présentent des traces de découpe réalisée à l'aide d'outils de pierre. On trouve des marques d'écorçage et de dépeçage (la peau et les chairs ont été retirées), ainsi que d'extraction du cerveau et de la moelle des os longs. Ces marques de découpe correspondent à celles que l'on observe sur les os des animaux retrouvés sur le site. Il s'agit de la plus ancienne preuve de cannibalisme humain connue à ce jour.

À la fin du XIX^e siècle, les restes de plus de 20 Néandertaliens inhumés dans les sables de la grotte de Krapina ont été découverts par le paléoanthropologue croate Dragutin Gorjanovic-Kramberger et, depuis, l'on s'interroge sur les Néandertaliens européens, qui vivaient entre 150 000 et 35 000 ans avant notre ère, étaient-ils cannibales ? Les os retrouvés brisés et éparpillés, portaient des marques de découpe. Malheureusement, ces os, fragiles, ont été extraits sans précaution et consolidés au moyen d'une épaisse couche d'une substance qui masque les éventuelles traces d'outils et qui complique les interprétations. Selon certains paléoanthropologues, les ossements des Néandertaliens de Krapina présentent des signes de cannibalisme. D'autres ont attribué les traces présentes sur les os à des pierres tombées du plafond de la grotte, à des morsures de carnivores (des loups, des hyènes...), ou à certains rites d'inhumation néandertaliens. De récentes analyses des os de Krapina et de Vindija, une autre grotte croate qui contient des restes de Néandertaliens et d'animaux plus récents, semblent indiquer que le cannibalisme aurait été pratiqué sur ces deux sites.

D'autres découvertes récentes confirment la pratique du cannibalisme chez certains hominidés. Sur les rives du Rhône, en Ardèche, Alban Defleur, de l'Université de Marseille, fouille depuis neuf ans la petite grotte de Moula-Guercy. Les Néandertaliens ont occupé ce site il y a environ 100 000 ans. Dans une des couches de sédiments, l'équipe a découvert les restes d'au moins six individus (un enfant de six ans et des adultes). A. Defleur a mené des fouilles méticuleuses, et il a établi des relevés précis, à la façon des experts médico-légaux qui enquêtent sur les lieux d'un crime. Il a relevé la position (selon les trois dimensions de l'espace) de chaque fragment osseux, animal ou néandertalien, de chaque indice botanique, de chaque outil de pierre. Grâce à ces travaux précis

A. Defleur a compris comment les os étaient éparpillés autour d'un foyer éteint depuis 100 000 ans.

En outre, les analyses microscopiques des fragments osseux de Néandertaliens et de la faune ont abouti aux mêmes conclusions que celles de l'équipe espagnole travaillant sur le site, plus ancien, de Gran Dolina. Certains Européens du Paléolithique ont donc pratiqué le cannibalisme.

La pratique cannibale semble même avoir été assez fréquente. On connaît un site européen très ancien avec des restes d'hominidés : ils étaient cannibales. Les deux sites croates ont été occupés par des Néandertaliens séparés par des centaines de générations, pourtant le cannibalisme était pratiqué par les uns et par les autres. Le site néandertalien français de Moula-Guercy présente lui aussi des traces de cannibalisme. À tel point que la plupart des anthropologues ne se demandent plus aujourd'hui si les premiers hominidés étaient cannibales, mais plutôt pourquoi ils étaient cannibales.

Cannibalisme en Amérique

De même, des découvertes faites depuis 30 ans sur des sites au Sud-Ouest des États-Unis ont modifié notre opinion sur la culture des Anasazi qui vivaient dans la région de Four Corners (de 100 avant notre ère à 1600 après notre ère), où ils cultivaient le maïs et construisaient des villages et des habitations troglodytes à flanc de falaise. Les vestiges archéologiques de ces populations sont particulièrement riches. Christy Turner, de l'Université de l'Arizona, a étudié des restes de squelettes humains brisés et brûlés retrouvés sur des sites anasazi d'Arizona, du Nouveau-Mexique et du Colorado. Site après site, il a retrouvé des signes de pratiques cannibales. Pourtant, rien dans l'histoire de la région ne laisse supposer que le cannibalisme était une pratique répandue. D'ailleurs, certaines tribus actuelles, qui se déclarent les descendants des Anasazi, sont embarrassées par la révélation du cannibalisme chez leurs ancêtres.

La majorité des sépultures anasazi présentent des squelettes en « connexion anatomique », accompagnés de céramiques décorées, très prisées des collectionneurs. Selon C. Turner, plusieurs dizaines de sites ont livré des restes humains fragmen-

és, souvent brûlés, qui évoquent une pratique plus courante que l'inhumation. Au cours des 30 dernières années, le nombre total d'ossements humains découverts sur ces sites se compte par dizaines de milliers. Ils correspondent à des individus répartis sur des dizaines de milliers de kilomètres carrés du Sud-Ouest des États-Unis, et sur plus de 800 années. Par exemple, il y a 10 ans, j'ai étudié un assemblage qui provenait d'un site anasazi du Canyon Mancos dans le Sud-Ouest du Colorado. Il était composé de 2 106 éléments osseux, issus d'au moins 29 hommes, femmes et enfants.

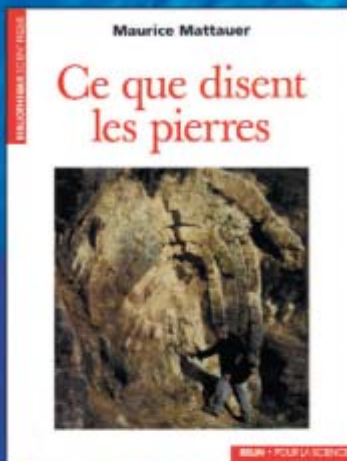
Ces assemblages d'ossements brisés ont été découverts dans des villages dont la taille variait de celle d'un petit hameau à celle d'une grande cité. Ils dataient souvent de la même époque que l'abandon des habitations. Ils présentent des traces de cuisson avant que la viande n'ait été prélevée. Le cerveau et la moelle des os longs étaient systématiquement prélevés, une fois la chair enlevée. Enfin, l'extrémité de certains morceaux d'os longs semble polie, comme s'ils avaient été cuits dans des récipients en céramique. Les fragments d'os humains de Mancos présentaient des marques de découpe qui correspondent à celles laissées par les mêmes Anasazi lorsqu'ils préparaient du cerb ou du mouton. D'après les données ostéologiques, les carcasses humaines étaient dépecées et rôties, débarrassées de leurs muscles et désarticulées. Les os longs étaient brisés sur des enclumes à l'aide de marteaux de pierre, les os spongieux écrasés et cuits dans des récipients de céramique. Ces résultats ont été controversés. Sans doute, pour de nombreux paléoanthropologues, le cannibalisme est une pratique tellement réprouvée que les preuves sont rejetées.

Il y a presque un an, Richard Marx et ses collègues de la Faculté de Médecine du Colorado ont publié la preuve la plus convaincante de cannibalisme connue sur les sites Anasazi. Ils ont fouillé trois puits, datés approximativement de 1150 de notre ère, sur le site de Cowboy Wash, près de Mesa Verde, dans le Colorado. Ils y ont relevé les mêmes éléments que sur les autres sites : les os humains étaient désarticulés, brisés, éparpillés, sans traces de sépultures. Grâce à leur excellente conservation, aux fouilles minutieuses, au prélèvement méticuleux

Histoires de cannibalisme

Divers récits ethnographiques de cannibalisme ont été consignés au fil des siècles dans diverses régions du monde. Certains sont des récits écrits par des témoins oculaires, d'autres sont des récits plus ou moins crédibles rapportés par des explorateurs, des missionnaires, des voyageurs ou des soldats. Ces deux vues d'artistes dépeignent le cannibalisme en Chine imposé par la famine, à la fin du XIX^e siècle, et dans le Nouveau Monde (d'après une gravure sur bois datant de 1497). De tels témoignages ne sont pas fiables.





CE QUE DISENT LES PIERRES

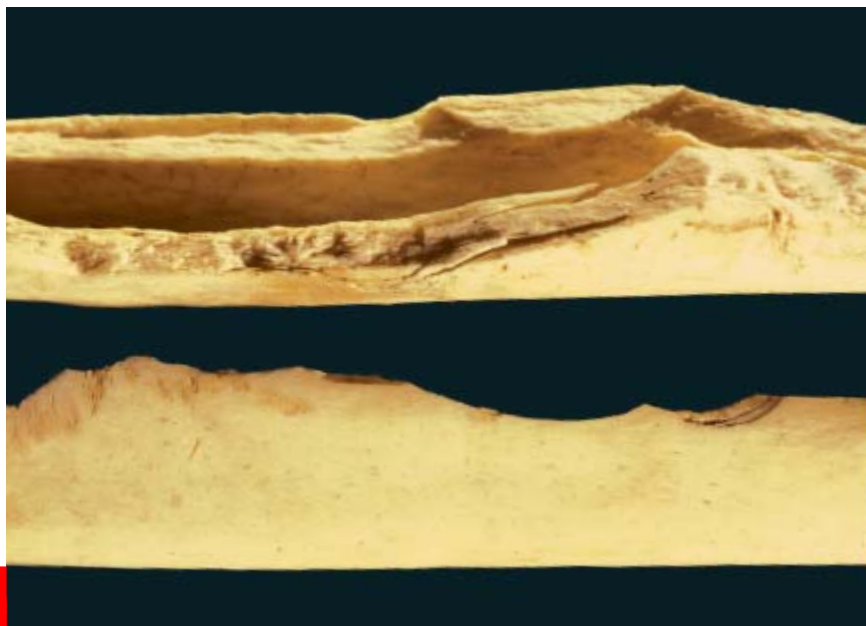
Maurice Mattauer

Ce livre vous convie à une promenade illustrée dans l'univers des pierres. À partir de l'image, Maurice Mattauer, professeur de géologie à l'Université de Montpellier, reconstitue l'histoire des roches et les grands événements dont elles sont les témoins, et invite à un «retour aux pierres», gage d'une nouvelle compréhension de la nature.

Code 075001
100 F • 15,24 €
144 pages

Voir bon de commande
page 96

BELIN • POUR LA SCIENCE



5. DES OS LONGS BRISES indiquent que l'on a prélevé la moelle. La viande (la graisse, les muscles et les autres tissus), qui entourait les os, n'était pas la seule partie consommée. Les crânes étaient brisés pour en prélever le cerveau, et la moelle était extraite des os longs. Ces deux os longs, découverts sur un site anasazi du Sud-Ouest du Colorado, ont été fendus sur toute leur longueur à l'aide de marteaux de pierre, afin d'accéder à la moelle.

des échantillons, et aux analyses chimiques du site, les chercheurs ont obtenu des preuves indiscutables de cannibalisme. Ils ont en effet découvert dans une poterie des restes de myoglobine humaine (une protéine présente dans le cœur et dans les muscles), ce qui indiquerait que de la chair humaine y a été cuite. Un coprothèque humain (un excrément fossilisé) non calciné, trouvé dans le foyer de l'une des habitations abandonnées, a également montré la présence de myoglobine humaine. D'après les données ostéologiques, archéologiques et biochimiques, de la chair humaine a bien été préparée et consommée à Cowboy Wash, ce qui conforte les hypothèses de cannibalisme préhis-

torique proposées pour les autres sites anasazi du Sud-Ouest des États-Unis.

Il n'en demeure pas moins qu'il est beaucoup plus difficile d'expliquer pourquoi le cannibalisme était pratiqué que de constater qu'il a existé. Les hommes mangent généralement parce qu'ils ont faim. Comment savoir si le cannibalisme représentait un moyen de subsister aux famines ou de se débarrasser des ennemis? Même dans le cas des Anasazi, qui ont fait l'objet de nombreuses études, il est impossible de savoir si le cannibalisme était lié à des pratiques religieuses, à la famine ou à une combinaison des deux, ou à d'autres raisons encore. Quelle qu'en ait été la cause, le cannibalisme fait désormais partie de notre passé collectif.

Jim WHITE est co-directeur du Laboratoire d'études de l'évolution humaine et professeur dans le département de biologie intégrée de l'Université de Berkeley en Californie.

P. VILLA et al, *Cannibalism in the Neolithic*, in *Science*, vol. 233, pp. 431-437, 1986.

P. VILLA, J. COURTIN, D. HELMER, P. SHIPMAN, C. BOUVILLE, E. MAHIEU, *In cas de cannibalisme au Néolithique*, in *Gallia Préhistoire*, vol. 29, pp. 431, 1986.

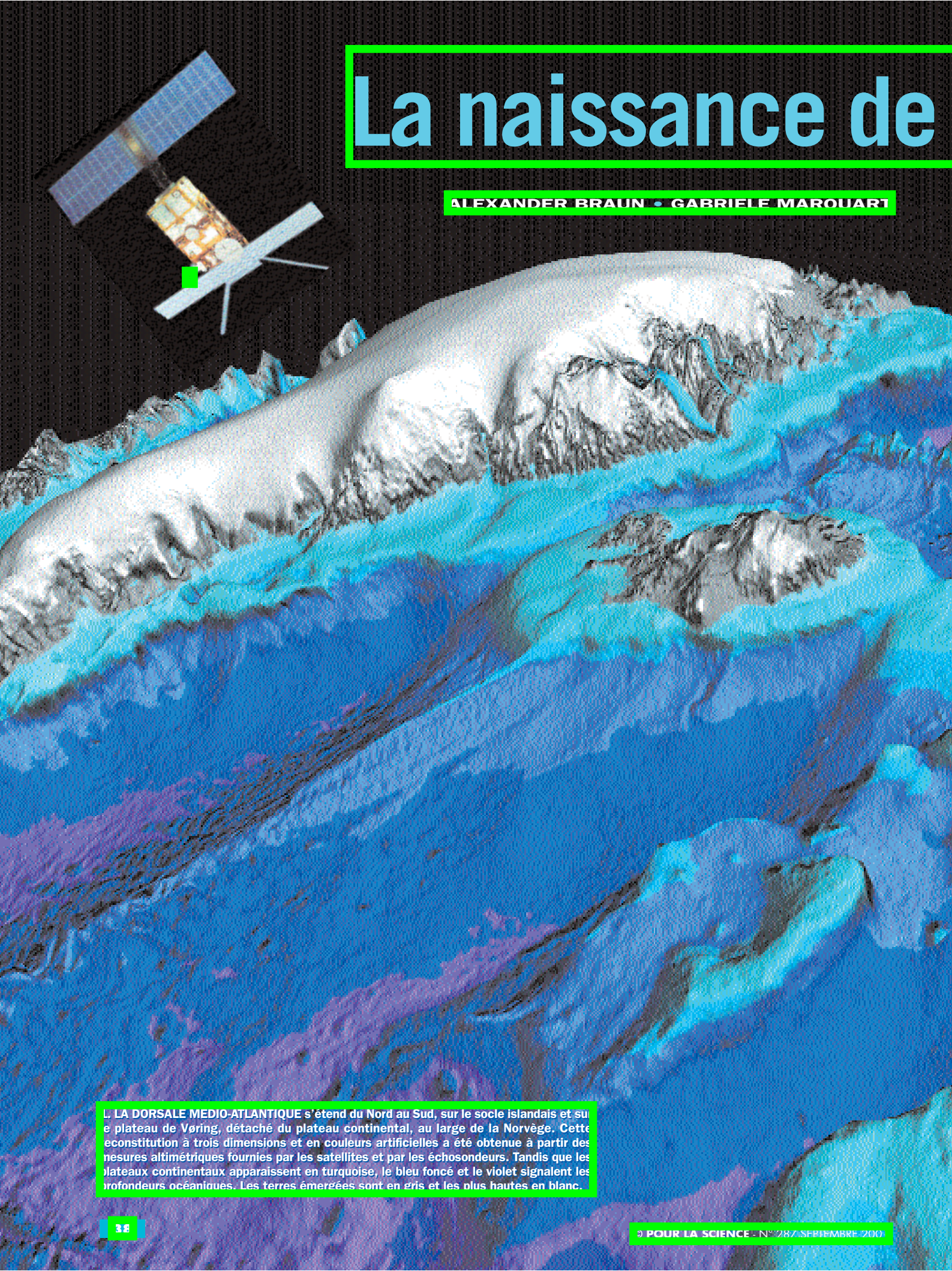
T. WHITE, *Prehistoric cannibalism at Mesa Verde 5MTUMR-2346*, Princeton University Press, 1992.

L. OSBORNE, *Does man eat man? Inside the great cannibalism controversy*, in *Lingua Franca*, vol. 7, N°4, pp. 28-38, 1997.

D. DEGUSTA, *Fijian cannibalism osteological evidence from Navatu*, in *American journal of physical anthropology*, vol. 110, pp. 215-241, 1999.

DEFLEUR et al, *Neanderthal cannibalism at Moula-Guercy, Ardèche, France*, in *Science*, vol. 286, pp. 128-131, 1999.

R. MARLAR et al, *Biochemical evidence of cannibalism at a prehistoric puebloan site in southwestern Colorado*, in *Nature*, vol. 407, pp. 74-78, 2000.

A 3D bathymetric map of the Mid-Atlantic Ridge and surrounding oceanic features. The map is color-coded: white and grey for land, light blue for continental shelves, and various shades of blue and purple for the deep ocean floor. A satellite is shown in the upper left corner, orbiting the Earth. The title 'La naissance de' is prominently displayed in a large, bold, black font at the top right.

La naissance de

ALEXANDER BRAUN • GABRIELE MAROUART

L LA DORSALE MEDIO-ATLANTIQUE s'étend du Nord au Sud, sur le socle islandais et sur le plateau de Vøring, détaché du plateau continental, au large de la Norvège. Cette reconstitution à trois dimensions et en couleurs artificielles a été obtenue à partir des mesures altimétriques fournies par les satellites et par les échosondeurs. Tandis que les plateaux continentaux apparaissent en turquoise, le bleu foncé et le violet signalent les profondeurs océaniques. Les terres émergées sont en gris et les plus hautes en blanc.

l'Atlantique Nord

Longtemps, la dorsale qui sillonne l'Atlantique a hésité, se faufilant d'abord entre le Groenland et le Labrador, avant de s'orienter vers le Nord-Ouest de l'Islande. L'altimétrie satellitaire et la tomographie sismique nous révèlent comment.

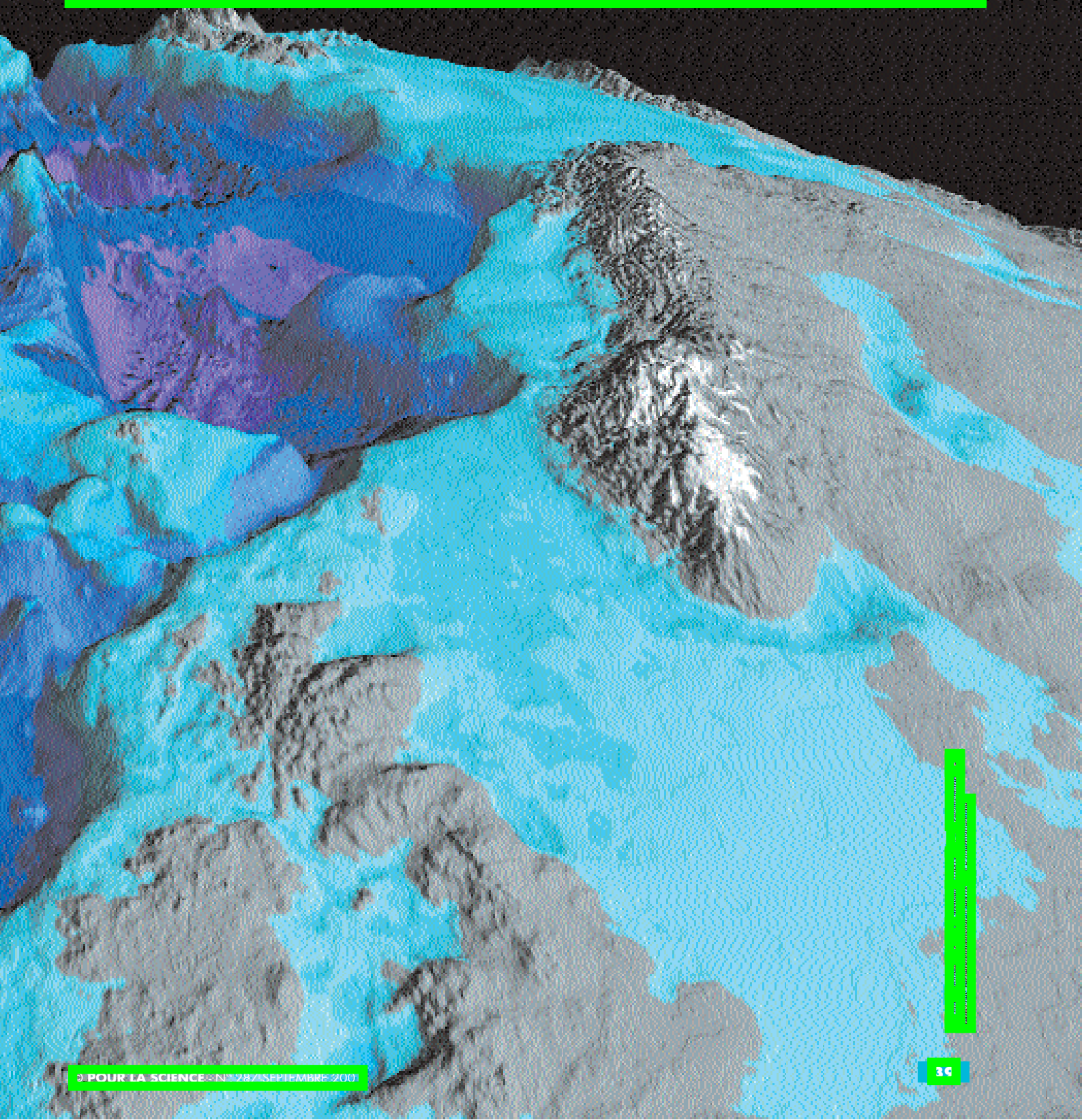


PHOTO: J. L. BOUILLON / IFREMER / OCEANUS
SCIENTIFIC DATA: J. L. BOUILLON / IFREMER / OCEANUS
SCIENTIFIC DATA: J. L. BOUILLON / IFREMER / OCEANUS

l'étonnante complémentarité des côtes de part et d'autre de l'Atlantique a été soulignée tant par le philosophe Francis Bacon en 1620

qu'en 1827 par le naturaliste Alexander von Humboldt, le père de la géographie physique. En 1596, le cartographe flamand Abraham Ortelius propose que «l'Amérique a été arrachée de l'Europe et de l'Afrique... par les tremblements de terre et des raz-de-marée». En 1756, le théologien allemand Theodor Lilienthal découvre la confirmation biblique de cette séparation (Genèse, chapitre 10, verset 25) : «Il naquit à Héber deux fils : le nom de l'un d'eux était Péleg, parce que de son temps la Terre fut partagée...» En 1858, la première ébauche d'explication rationnelle de la complémentarité des côtes d'Europe et d'Amérique est proposée. Impressionné par la ressemblance des flores fossiles d'Europe et d'Amérique du Nord au carbonifère (300 millions d'années avant notre ère), le géographe Antonio Snider-Pellegrini émet l'hypothèse que les continents se sont déplacés à la surface de la Terre. L'Américain est toutefois si dévôt qu'il attribue le phénomène au... Déluge!

C'est surtout Alfred Wegener, l'inventeur de la dérive des continents qui a fait progresser l'étude des fractures continentales. Dans sa théorie, publiée en 1912, il énumère de nombreux arguments géologiques en faveur d'une origine commune des masses continentales situées de part et d'autre de l'Atlantique. Toutefois, il n'explique pas de façon convaincante la migration supposée des continents. L'explication du phénomène ne voit le jour qu'à la fin des années 1960, après la découverte de l'expansion des fonds océaniques : la théorie de la tectonique des plaques se développe alors.

Selon cette théorie, l'enveloppe la plus externe de la Terre, la lithosphère, est divisée en plaques, rigides qui «flottent» sur une couche plus molle, l'asthénosphère. D'une épaisseur moyenne d'environ 100 kilomètres, elles sont constituées de croûte terrestre et d'une partie du manteau supérieur (couche constituée des 660 premiers kilomètres sous l'écorce terrestre). Les plaques s'écartent les unes des autres, repoussées par le magma qui remonte des profondeurs avant de se refroidir et de se solidifier.

Les zones de plancher océanique qui se forment ainsi de part et d'autre d'une dorsale croissent régulièrement. Elles acquièrent une aimantation sous l'influence du champ magnétique terrestre, lequel s'inverse à des intervalles compris entre 100 000 ans et un million d'années. Ainsi, on observe une alternance de bandes d'aimantations opposées à la surface du plancher océanique. Si, au centre de l'océan, les plaques croissent et s'écartent, elles subissent, dans des zones dites «de subduction», qui sont souvent proches des continents, une compression qui les fait plonger dans le manteau terrestre, où elles fondent. Des abysses, c'est-à-dire de profondes fosses océaniques, sont souvent présentes au voisinage des zones de subduction.

Les continents constituent les reliefs les plus élevés de certaines de ces plaques et se déplacent avec elles. Toutefois, contrairement au plancher océanique, ils ne peuvent pas plonger dans le manteau au niveau des zones de subduction. Ainsi, lorsque deux continents entrent en collision, ils se soudent l'un à l'autre, comme c'est par exemple le cas de l'Inde et de l'Asie ; au contraire, les masses continentales peuvent aussi se fissurer et se séparer, s'écartant progressivement. La mer s'engouffre alors entre les deux, et le bras de mer qui se forme finit par devenir un océan. Le Nouveau Monde et l'Ancien Monde se sont séparés de cette façon.

Les géologues sont confrontés pour cette raison à deux types de croûte terrestre : la croûte océanique et la croûte continentale. L'une et l'autre sont difficiles à étudier, mais pour des raisons différentes. Nous avons déjà évoqué, la croûte océanique s'enfonce avec le manteau lithosphérique (la partie supérieure du manteau) vers l'intérieur de la Terre au niveau des zones de subduction, tandis que la croûte continentale, plus légère, reste en surface (pour la plus grande partie). Très anciens, les continents conservent les traces de leur évolution tectonique depuis plusieurs centaines de millions d'années, voire plusieurs milliards d'années (la formation de la croûte terrestre a commencé il y a quatre milliards d'années environ). Les informations géologiques s'y sont accumulées, recouvertes, entremêlées... Au contraire, la croûte océanique est jeune, souvent moins de 180 millions d'années. Seuls les événements tectoniques récents



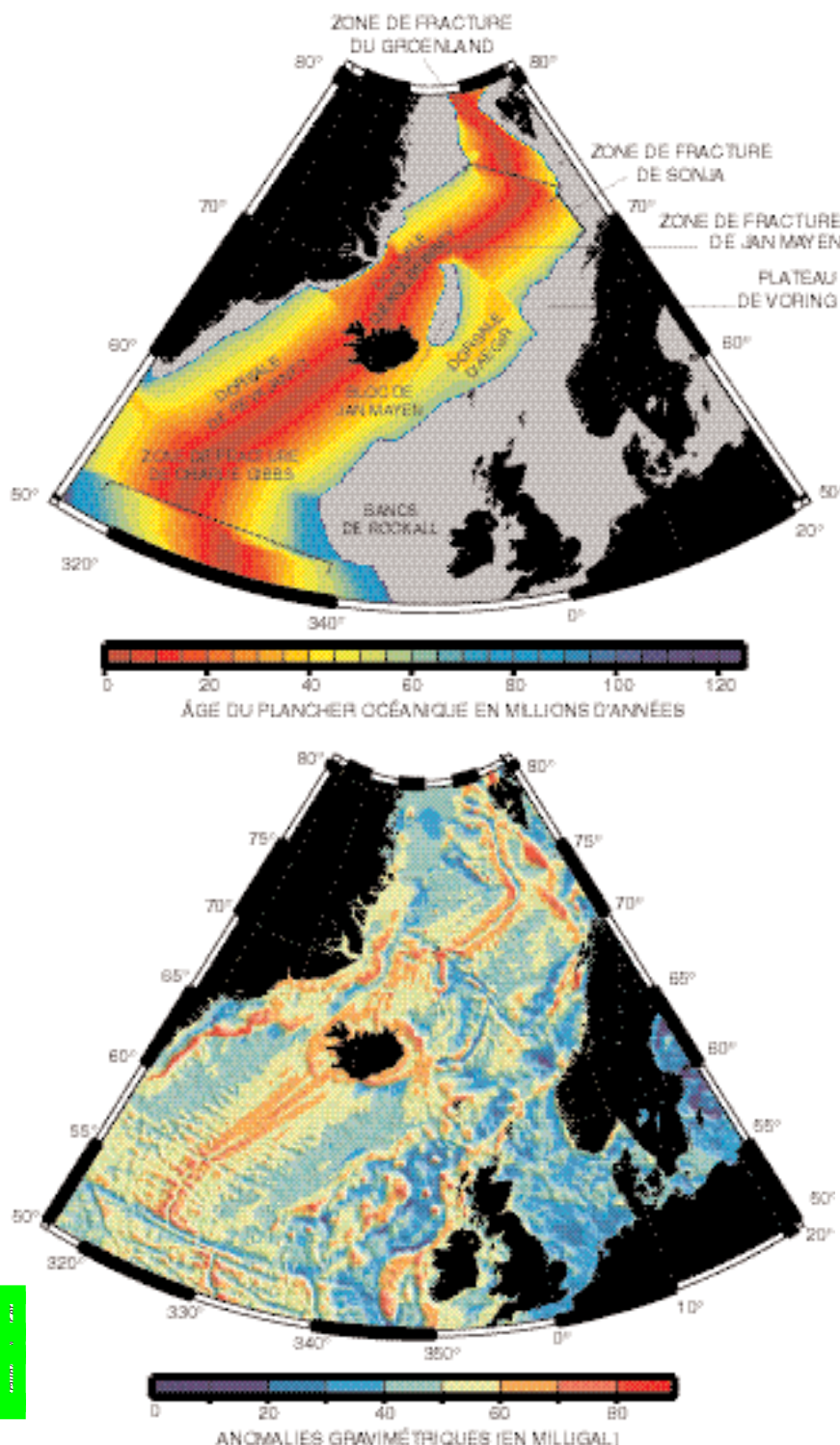
2. L'ATLANTIQUE NORD qui s'étend, pour les auteurs, entre le 52° parallèle (au Sud de l'Irlande) et 80° parallèle (au Nord du Spitzberg)

sont imprimés à sa surface ; ils racontent l'évolution « récente » de l'océan. Les traces des événements anciens ont plongé dans le manteau avec les plaques anciennes. En fait, si l'étude de la croûte océanique est difficile, c'est parce que d'épaisses couches de sédiments et d'eau la recouvrent.

L'histoire de l'Atlantique

Les premières reconstructions de l'histoire de l'Atlantique reposaient sur des mesures de la profondeur de l'océan par échosondeur (c'est-à-dire par réflexion d'ultrasons) et sur quelques rares forages en haute mer. Puis, les données bathymétriques et la mesure de la profondeur des mers dévoilèrent la topographie du plancher atlantique. Les résultats, grossiers, avaient pourtant révélé la présence de la dorsale médio-atlantique, qui traverse l'océan du Nord au Sud. Chaîne de montagnes sous-marines, la dorsale médio-atlantique est coupée, de proche en proche, par des zones de fractures, qui sont autant d'entailles profondes, où deux sections de dorsales peuvent être décalées l'une par rapport à l'autre de plusieurs centaines de kilomètres. Les forages ont aussi livré des informations sur les fossiles caractéristiques des sédiments des diverses époques et sur l'aimantation de la croûte basaltique sous-jacente, dont on déduit l'âge des roches étudiées.

Aujourd'hui, de nouvelles techniques d'observation révèlent les mécanismes géodynamiques de formation de la lithosphère. La mesure du niveau des mers par satellite (altimétrie satellitaire) est l'une de ces méthodes ; la tomographie sismique – on se sert des ondes émises lors de tremblements de terre pour observer l'intérieur de la Terre – en est une autre. Ces nouvelles méthodes ont fait évoluer la « vision classique » de la formation de l'Atlantique, telle qu'elle était envisagée jusqu'à récemment par la tectonique des plaques. À l'Université de Francfort, à l'aide des nouvelles méthodes, nous avons étudié la portion de l'Atlantique qui s'étend du 52° au 80° parallèle (du Sud de l'Irlande au Nord de la Norvège), que nous nommerons « Atlantique Nord » dans la suite. Nous proposons ici une esquisse globale de la formation de cette portion d'océan, placée dans le contexte général de celle de l'Atlantique.



3. GRACE À LA CARTOGRAPHIE MAGNÉTIQUE du plancher océanique et à l'altimétrie par satellite, les géologues reconstituent l'histoire de l'Atlantique Nord. Entre 1994 et 1995, le satellite européen ERS-1 a livré plus de 500 000 mesures qui leur ont permis de tracer la carte des variations du champ de gravitation (*en bas*). Comme les variations gravimétriques sont dues pour l'essentiel au relief, leur mesure livre une image assez fidèle du fond de l'océan. Ainsi, des écarts positifs par rapport à la moyenne (*en rouge et en jaune*) traduisent la présence de montagnes ou de volcans sous-marins le long de la dorsale médio-atlantique. Inversement, les fosses océaniques ou les bassins comblés de sédiments, dont la densité est inférieure à celle de la roche magmatique apparaissent en bleu. La dorsale d'Aegir, ainsi que les anomalies alternativement positives et négatives qui s'étirent sur le plateau continental des Lofoten (*voir page de gauche*) sont difficiles à identifier sur les cartes topographiques. L'âge du plancher océanique augmente régulièrement des deux côtés de la dorsale médio-atlantique (*en haut*). Ce plancher est alimenté en continu par du magma, qui, en refroidissant, forme la croûte océanique. Les roches crustales superficielles conservent l'aimantation acquise lors du refroidissement, ce qui reflète les inversions du champ magnétique terrestre et permet ainsi de dater les différentes zones du plancher océanique. Sur cette illustration, les zones grises correspondent aux plateaux continentaux.

Nous avons choisi d'étudier l'Atlantique Nord, parce qu'il se prête à une analyse géologique poussée. De nombreuses données sont déjà disponibles, et l'évolution tectonique de cette zone a en plus l'intérêt d'être liée aux mouvements profonds qui ont lieu dans le manteau ; leur importance est attestée par le volcanisme actif qui y règne, que ce soit sous l'eau ou sur terre (en Islande, sur l'île de Jan Mayen...).

Les prémisses

L'ouverture de l'Atlantique Nord a été précédée par une série de violentes éruptions, il y a quelque 56 millions d'années. Une zone de fracture de quelque 3 000 kilomètres de longueur se forme alors, du Sud de l'Angleterre au Nord de la Norvège. La mer l'emphahit, puis en deux millions d'années seulement, une masse de plus de deux millions de kilomètres cubes de laves basaltiques se déverse à la surface. Les restes de cet énorme apport constituent aujourd'hui des bandes symétriques de basaltes de plusieurs kilomètres d'épaisseur qui s'étirent à l'Est du Groenland et au Nord-Ouest du continent européen. La dorsale

océanique qui se forme ensuite au fond de ce rift continental a un comportement relativement normal, si l'on met à part l'apparition de quelques centres volcaniques actifs, qui engendrent le plateau norvégien externe ou plateau de Vøring, reliant les îles Féroé à l'Islande, ces îles elles-mêmes ainsi que les îles Shetland (voir la figure 2).

Il y a environ 38 millions d'années une nouvelle phase volcanique intense se déclenche. Elle est particulièrement violente à proximité de l'Islande. Depuis, l'île et son socle constituent l'une des régions volcaniques les plus actives de la Terre. Plus d'un million de kilomètres cubes de matériaux volcaniques se sont accumulés dans la région. Un tel débit indique que des matières chaudes remontent sous l'Islande des grandes profondeurs, voire du manteau inférieur (2 880 kilomètres de profondeur). Ce type de flux vertical qui s'élargit vers le haut est nommé « panache ». Lorsque les roches mantelliques, que le panache transporte, parviennent à une centaine de kilomètres sous la surface, la pression diminue ce qui entraîne leur fusion partielle. C'est ainsi que se forme la lave qui alimente les volcans.

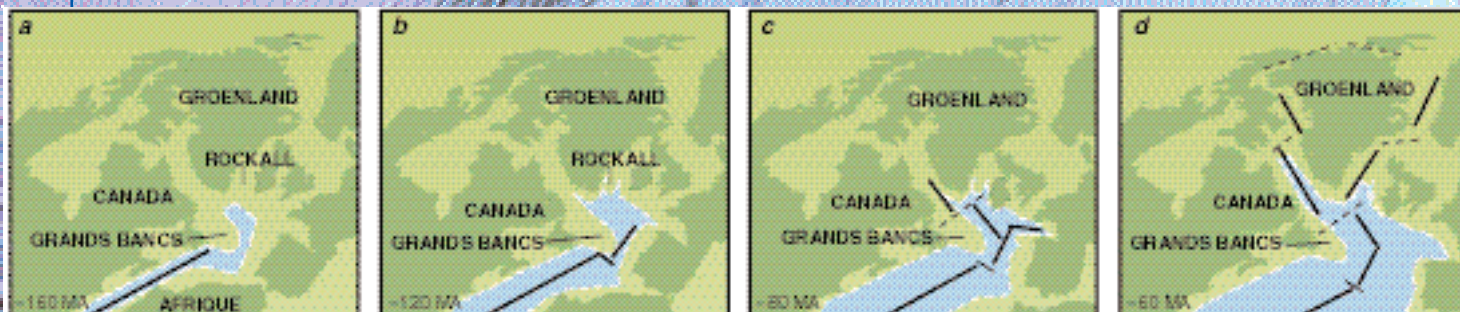
Un tel panache semble avoir été aussi à l'origine de la séparation des masses continentales européennes et américaines, il y a 56 millions d'années. Les représentations de l'évolution de cette région ont longtemps été simplistes : elles indiquaient que, dès sa naissance, la dorsale médio-atlantique s'étendit vers le Nord à partir de la zone de fracture de Charlie Gibbs (aujourd'hui entre l'Islande et la presqu'île du Labrador, située au Nord-Ouest du Québec). Un important réservoir de magma, sous l'Islande, aurait favorisé la formation de la croûte océanique. La montée continue de ce magma aurait écarté la croûte nouvellement formée, élargissant ainsi l'océan. La topographie du plancher océanique et la façon dont les contraintes et la répartition des températures à l'intérieur de la Terre façonnaient le panache ont été négligées jusqu'à présent. Les anciens modèles de la formation de l'Atlantique Nord ont été élaborés en ne tenant compte que des structures les plus importantes : les plaques, la dorsale médio-atlantique, et les zones de fracture. Jusqu'à présent, on décrivait l'ouverture de l'Atlantique Nord sans tenir compte de la dynamique des phénomènes géologiques.

La formation de l'Atlantique Nord

Il y a 500 millions d'années, l'Europe est séparée du bloc constitué de l'Amérique du Nord et du Groenland par un océan nommé Iapetus. Cet océan primitif commence à se fermer vers -450 millions d'années, et le bloc de l'Amérique du Nord et du Groenland se rapproche alors de l'Europe. Il y a 400 millions d'années, l'Amérique du Nord et l'Europe entrent en collision, ce qui conduit à la formation de la chaîne calédonienne, dont les vestiges constituent aujourd'hui encore les hauteurs de la Norvège, de l'Écosse et de certaines régions de l'Est des États-Unis. De -400 à -290 millions d'années, les masses continentales du monde entier se réunissent en un supercontinent, la Pangée. Avant même que cette unification ne soit achevée, de nouveaux fossés d'effondrement (ou rifts) réapparaissent entre la Scandinavie et le Groenland. Il y a 220 millions d'années, ce système de rifts se prolonge jusqu'au banc de Rockall, à l'Ouest de l'Écosse, et forme une mer peu profonde. Ce n'est qu'il y a 160 mil-

lions d'années, qu'une mer plus profonde – l'Atlantique – commence à s'ouvrir entre l'Est des États-Unis et l'Ouest de l'Afrique.

L'Atlantique sépare d'abord l'Afrique de l'Ouest de l'Amérique du Nord (a ; sur les différents schémas, on a indiqué les dorsales par des lignes noires). Au cours de la phase suivante (b), un bras de l'Atlantique s'ouvre entre Terre-Neuve et l'Espagne et, simultanément, la croûte continentale s'étire sur une large zone à proximité du banc de Rockall (du nom d'une petite île britannique située au Nord de l'Irlande) entre le Groenland et l'Europe du Nord. Il y a environ 80 millions d'années, deux nouvelles branches de la dorsale médio-atlantique se forment (c) ; tandis que l'une s'étend vers le Nord-est dans le golfe de Gascogne, repoussant la presqu'île ibérique vers le Sud, l'autre s'étend vers le Nord-Ouest jusqu'au banc de Rockall. Alors que l'activité de cette première dorsale s'arrête rapidement, la dorsale Nord-Ouest se prolonge entre le Groenland et l'



Pour améliorer leur analyse de la formation de l'Atlantique Nord, les géologues devaient, d'une part, cartographier le plus précisément possible les structures géologiques et tectoniques présentes au fond de l'océan et, d'autre part, analyser les mouvements dans le manteau. Pour cela, il fallait repérer les variations de température et de densité à l'origine des courants de convection qui parcourent le manteau profond.

Satellites et sismomètres

Il y a une quinzaine d'années, ce progrès devint possible en raison des nouvelles méthodes de mesure disponibles. Dans les années 1980, l'altimétrie satellitaire et la tomographie sismique s'imposent dans l'étude géologique des fonds marins. Ces deux techniques de mesure complètent les méthodes anciennes, telles que la mesure bathymétrique par échosondeur et les forages en mer. À l'aide de signaux radar, les satellites mesurent en permanence leur altitude, c'est-à-dire leur hauteur au-dessus de l'océan. Toutefois, la surface de la mer est rarement plane, puisqu'elle est poussée par les vents et agitée de vagues. Par ailleurs,

le niveau moyen de la mer en un endroit dépend de l'attraction gravitationnelle des masses qui constituent localement le Globe. Cette attraction est modulée par la topographie et les densités des roches au fond de la mer. C'est pourquoi la structure de la distribution des masses sur le plancher océanique se « décalque » à la surface des mers.

L'ensemble des données enregistrées par les satellites ERS-1 et ERS-2 lancés en 1995 par l'Agence spatiale européenne, a fourni, par cette méthode, une carte précise des fonds océaniques. Le satellite de télédétection ERS-2 couvre la surface de la Terre en 35 jours. Il parcourt quelque 20 millions de kilomètres et procède à une mesure tous les 300 mètres ! Avec un bateau, il faudrait plus de 60 ans pour enregistrer autant de données. Les mesures faites à bord des navires n'en restent pas moins importantes, ne serait-ce que pour vérifier et étalonner les données satellitaires.

Une fois le fond océanique connu par des images, il reste à comprendre quels phénomènes géologiques l'ont formé. La tomographie sismique fournit des images de l'intérieur de la Terre. Elle procède du même principe que

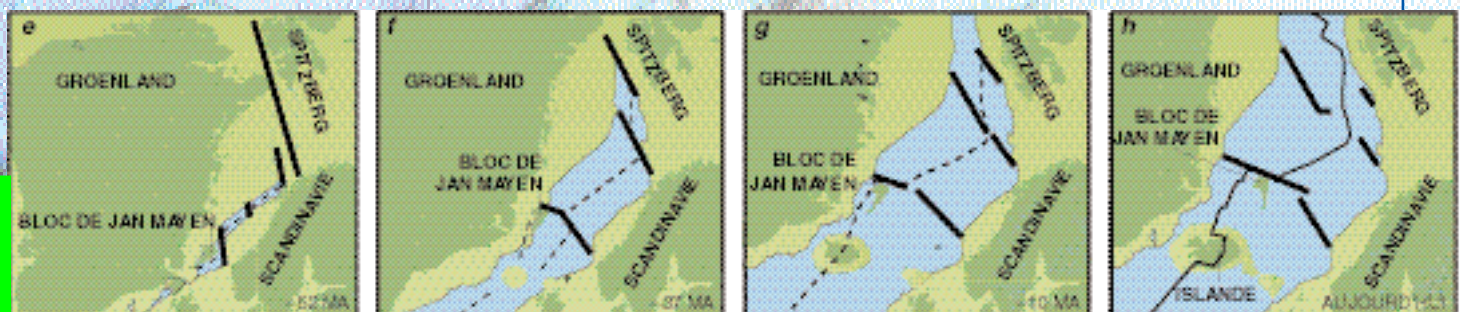
la tomodensitométrie, mise en œuvre dans les scanners d'hôpitaux pour visualiser l'intérieur du corps humain. Toutefois, alors qu'en médecine, on utilise des rayons X ou des ultrasons pour « éclairer » l'intérieur de l'organisme, en sismologie on se sert des ondes sismiques émises par les tremblements de terre. Dans les stations de surveillance sismologique, on enregistre le plus grand nombre possible de séismes et, lors de chaque tremblement de terre, on mesure le laps de temps que mettent les ondes sismiques pour se propager depuis l'épicentre ; ces ondes n'ont pas toutes la même vitesse, laquelle dépend de la densité des roches qui constituent le manteau terrestre. Ces variations de densité résultent de changements de température. Grâce à la tomographie sismique, les géologues cartographient les zones chaudes et froides du manteau terrestre. Comme la densité des roches diminue et que leur ascension dans le manteau s'accélère lorsque la température augmente, les régions les plus chaudes du manteau – les points chauds – constituent des zones où des roches visqueuses remontent ; à l'inverse, dans les régions plus froides, les roches s'enfoncent.

Canada (d). L'Atlantique poursuit son expansion vers le Nord entre le Nord-Est du Canada et le Groenland, ce qui donne naissance à la mer du Labrador.

Il y a quelque 60 millions d'années, une énorme masse de roches particulièrement chaudes remonte du manteau terrestre sous le Groenland et sous l'actuelle Islande. À mesure que la pression diminue, ces roches fondent partiellement, produisant de grandes quantités de laves. L'Atlantique se prolonge vers le Nord en formant une seconde branche entre le Groenland et le bloc constitué de l'Écosse et de la Scandinavie (d). Une zone de fracture de quelque 3 000 kilomètres de longueur apparaît entre le Sud de l'Angleterre et le Nord de la Norvège. La plaque européenne se serait alors séparée de la plaque Nord-américaine, à l'Est du Groenland. Le Canada et le Groenland auraient alors cessé de s'écarter l'un de l'autre.

Il y a entre 56 et 52 millions d'années, l'Atlantique Nord commence à s'ouvrir jusqu'à la zone de fracture dite de Senja-Groenland, qui s'étend de la côte Nord-Ouest de la Norvège au bassin arctique (e).

De nouvelles zones de fracture apparaissent sans cesse et modifient la position de la dorsale (en pointillés). Il y a 37 millions d'années (f), la croûte océanique cesse de se former dans la mer du Labrador, à l'Ouest du Groenland. La zone de fracture au Sud du Spitzberg se transforme alors progressivement en un système de dorsales, qui se prolonge jusque dans le bassin arctique. Une intense activité volcanique apparaît en Islande. Au Nord de l'île, la dorsale se sépare en deux branches. Entre -37 millions d'années et -10 millions d'années, sous l'action de la branche occidentale de cette dorsale, le bloc continental de Jan Mayen se détache, tandis que la branche orientale – la dorsale d'Aegir – s'éteint. L'Islande a alors presque sa taille actuelle et devient partie intégrante du système de la dorsale médio-atlantique (g). Aujourd'hui, le plancher océanique s'ouvre de un à deux centimètres par an au niveau des dorsales dans presque tout l'Atlantique Nord, tandis qu'il continue à s'épaissir par accréation de matériau qui remonte du manteau. La ligne qui sépare l'Europe de l'Amérique du Nord passe désormais entre la Scandinavie et le Groenland.



Un plancher océanique révélateur

Ces méthodes ont changé la conception qu'avaient les géologues du plancher océanique de l'Atlantique Nord et de son origine. De fait, le fond du grand océan qui baigne l'Europe est désormais mieux connu par endroits que certaines régions continentales solées ! La carte du fond de l'Atlantique Nord obtenue à partir des données du satellite ERS-1 (voir la figure 3) révèle une profusion de structures qui « racontent » le déroulement de la formation de l'océan (voir l'encadré de la page 42). La structure la plus importante est le système de dorsale médio-atlantique. Elle court depuis la zone de fracture de Charlie Gibbs, passée par la dorsale de Reykjanes (au Sud de l'Islande) et rejoint le socle islandais par le Sud.

En Islande, cette dorsale se ramifie. Sa branche occidentale se raccorde au Nord de l'île, à la dorsale de Kolbeinsey qui constitue la zone principale du système, à proximité du bord continental du Groenland. Des vestiges d'une dorsale plus ancienne – la dorsale d'Aegir – sont localisés à l'Est du Groenland ; cette dorsale a cessé d'être active il y a 27 millions d'années environ, avant que le centre d'activité volcanique de cette région ne se déplace vers l'Ouest pour rejoindre la dorsale de Kolbeinsey (probablement en deux étapes). Le centre volcanique de la dorsale de Kolbeinsey et les vestiges de celui qui l'a précédé à l'Est sont limités au Nord par la zone de fracture de Jan Mayen (voir les figures 2 et 3).

L'île de Jan Mayen (372 kilomètres carrés) est une île volcanique norvégienne située au Nord-Ouest de l'Islande. Ce fragment continental serait détaché du plateau continental groenlandais, lorsque l'axe de la dorsale s'est déplacé vers l'Ouest, à cet endroit. La dorsale, aujourd'hui active, s'étend au-delà de l'Islande vers le Nord, puis vers le Nord-Ouest sans rencontrer d'obstacles. Dans le passé, une zone de fracture longeait le Nord du Groenland et coupait ce qui allait devenir l'extrémité Nord de la dorsale médio-atlantique. Il y a quelque 37 millions d'années, les plaques océaniques se sont écartées, laissant remonter du magma et de la croûte océanique s'y est accumulée.

L'image tomographique du manteau terrestre sous l'Atlantique Nord

(voir la figure 5) fait apparaître deux zones où la vitesse de propagation des ondes sismiques est réduite. La vitesse de propagation des ondes sismiques étant inversement proportionnelle à la racine carrée de la densité des roches, on en déduit que dans cette région, la densité est faible et que la température est supérieure à celle des zones environnantes, en raison de la remontée de matière issue du manteau. La première de ces zones de remontée des roches mantelliques se trouve au Sud du Groenland et a une forme tubulaire ; elle prend naissance à quelque 2 880 kilomètres de profondeur, à la limite entre le manteau et le noyau liquide, et remonte en direction du Nord. Elle atteint le manteau supérieur à 660 kilomètres de profondeur sous le plateau continental européen, à la latitude de la Grande-Bretagne. Elle semble se ramifier en plusieurs branches, dont l'une rejoint la dorsale médio-atlantique au niveau de l'Islande.

La seconde zone, particulièrement chaude, qui apparaît sur l'image tomographique ne semble pas être reliée au manteau terrestre profond. Elle se situe sous le Groenland à une profondeur comprise entre 200 et 400 kilomètres seulement. Cette zone chaude suit une région où la lithosphère s'amincit en direction de l'Atlantique Nord jusqu'à la dorsale médio-atlantique, le long du plateau continental.

Les images tomographiques confirment que les zones où de la matière chaude remonte du manteau ont participé à l'ouverture de l'Atlantique Nord. Les volcans qui sont apparus tout le long de la dorsale ont d'abord été alimentés par la matière issue du manteau profond. Aujourd'hui, seuls continuent à fonctionner de cette façon les volcans islandais et la dorsale de Reykjanes au Sud de l'Islande ; l'intensité de ce volcanisme tend à décroître. Au Nord de l'Islande, le volcanisme est sans doute lié à l'existence d'une région chaude sous le Groenland.

Si un puissant panache situé sous l'Islande a formé l'Atlantique Nord que nous connaissons, que s'était-il passé avant ? La « préhistoire » de la région remonte au Cambrien, il y a quelque 500 millions d'années. À cette époque, l'Amérique du Nord et l'Europe étaient séparées par un océan étroit nommé « Iapetus » (l'un des titans de la mythologie grecque). Ce pré-

cursor de l'actuel Atlantique Nord commença à se fermer il y a 450 millions d'années, sous l'influence d'une zone de subduction. Il avait presque entièrement disparu vers -400 millions d'années, de sorte que l'Europe et l'Amérique du Nord entrèrent en collision. Comme la collision entre l'Inde et l'Asie a engendré l'Himalaya, cette collision conduisit à la formation d'une grande chaîne de montagnes : la chaîne calédonienne. Ses restes érodés sont encore visibles aujourd'hui en Norvège, en Écosse, au Groenland et dans l'Est des États-Unis. Puis, toutes les masses continentales se rassemblèrent en un supercontinent, la Pangée.

Préhistoire d'une séparation

Avec ses dimensions démesurées, la Pangée agissait comme une immense couche isolante, déposée à la surface de la Terre, et qui empêchait tout transfert de chaleur des profondeurs vers la surface. Les matériaux chauds s'accumulèrent dessous, ramollirent la lithosphère continentale et l'amincirent. Cette érosion thermique aboutit à la formation de zones d'expansion et de failles, qui amorcèrent l'ouverture de la masse continentale et la formation de nouveaux continents.

Les observations que l'on peut faire sur le bord du plateau continental norvégien et du plateau de Barents (en face du Nord de la Norvège et de la péninsule de Kola) montrent que ces mouvements d'expansion commencèrent dès la fermeture de l'océan Iapetus, avant même que la Pangée ne soit totalement soudée. Lors de cette première expansion, une large zone de volcans peu actifs et de fossés d'effondrement, nommés rifts, apparaît à la surface de la Pangée et s'étend progressivement vers le Sud. Vers -220 millions d'années, ce système de rifts atteint le banc de Rockall, aujourd'hui à l'Ouest de l'Écosse, et un système de failles apparaît dans la région de la mer du Nord. Remplies d'eau, ces failles forment des mers continentales peu profondes et dépourvues de croûte océanique. Les forages et l'étude des anomalies magnétiques des fonds marins ont révélé que l'Atlantique Nord a commencé à se constituer il y a seulement 160 millions d'années, entre ce qui est aujourd'hui l'Est des États-Unis et l'Ouest de l'Afrique. La dorsale médio-atlantique qui apparaît alors est limitée au Nord-Est par une zone de fracture à la latitude de ce qui est aujourd'hui le

L'altimétrie par satellite

L'altimétrie satellitaire est une technique de mesure du niveau de la mer à l'aplomb du satellite. Le satellite envoie un signal radar vers la surface de la mer et capte le signal renvoyé. Une horloge de haute précision mesure le temps d'aller-retour, dont un circuit électronique déduit la distance du satellite à la surface de l'océan.

Ce principe paraît simple, mais sa mise en œuvre est complexe. Si les ondes radar traversent l'atmosphère, quelles que soient les turbulences qui y règnent, leur vitesse de propagation varie en fonction de facteurs tels que l'humidité de l'air. Ainsi, les conditions météorologiques durant les mesures doivent être prises en compte lors de l'interprétation des résultats.

De plus, la position du satellite doit être connue précisément à chaque instant de la mesure. En raison des irrégularités du champ de gravitation terrestre, l'orbite d'un satellite est irrégulière. Le globe terrestre n'est pas, lui-même, une sphère parfaitement homogène : si sa surface est globalement ronde, elle présente des aspérités (montagnes) ; en outre, les masses sont irrégulièrement réparties au sein du Globe. Certains satellites sont dotés de dispositifs embarqués qui, en liaison avec des stations au sol, déterminent en permanence leur position ; d'autres sont localisés à partir du sol à l'aide d'impulsions laser.

L'inégale répartition des masses sous le plancher océanique résulte de mécanismes géologiques qui façonnent le relief et des puissants courants de convection qui brassent les roches dans le manteau. Les deux effets se conjuguent pour conférer à la surface moyenne de l'océan des hauteurs, c'est-à-dire des distances au centre de la Terre, différentes. Certains creux atteignent 10 mètres au Sud de l'Inde, et certaines bosses 90 mètres au large de la Nouvelle-Guinée.

Après le succès de premières expériences menées à bord de la station spatiale américaine *Skylab*, en 1973, toute une série de satellites furent dédiés à l'altimétrie satellitaire : *GEOS-3* en 1975, *SEASAT* en 1978, *GEOSAT* en 1985, puis *ERS-1* en 1991, *Topex-Poseidon* en 1992 et enfin *ERS-2* en 1995. La mise sur orbite de deux nouveaux altimètres est prévue pour 2001 : *Jason-1* et *ENVISAT*.

L'objectif des premières missions était la mesure du champ gravitationnel sous-marin sur l'ensemble du Globe, et l'obtention d'une image du plancher océanique. En effet, l'intensité du champ de gravitation augmente au-dessus des montagnes sous-marines, de sorte que l'eau s'y amoncelle. En revanche, la surface de l'océan se creuse à la verticale des zones où le champ de gravitation diminue. Même si ces effets ne sont pas très marqués, ils expliquent que la topographie du fond océanique se « décalque » à la surface de l'eau.

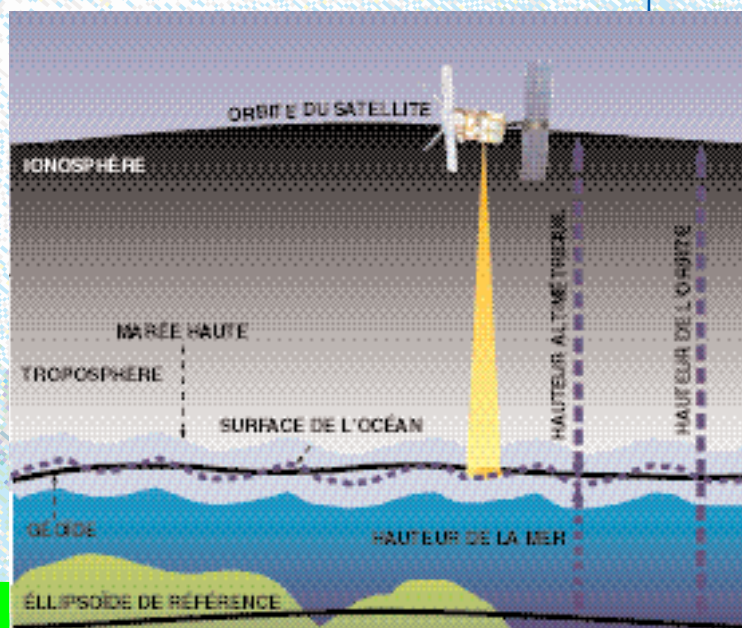
Aujourd'hui, grâce à l'altimétrie satellitaire, le fond marin a été cartographié sur la quasi-totalité de la planète. Il s'agit maintenant d'étudier précisément l'influence des systèmes de courants océaniques et des phénomènes océanographiques qui influent – eux aussi – sur la surface de l'eau. L'étude des régions de la Terre qui sont recouvertes de glace en permanence commence tout juste.

Suivant la tâche qui leur est dévolue, les satellites ont des orbites différentes. Leurs trajectoires passent près des pôles pendant que la Terre tourne, de sorte qu'ils survolent toute la Terre en spirale (à l'exception des latitudes les plus élevées). À l'issue d'un cycle, un satellite revient à sa position d'origine, de sorte que la durée du cycle varie en fonction de l'orientation et de la forme de la trajectoire. Lorsque les cycles sont longs, les données enregistrées sont denses, mais deux mesures effectuées au même endroit sont éloignées dans le temps. Quand les cycles sont courts, on enregistre peu de mesures dans un cycle, mais deux mesures

au même endroit sont proches dans le temps. Ainsi, un satellite qui repasse souvent au même endroit sert, par exemple, à étudier des phénomènes qui évoluent rapidement, tels les courants océaniques. En revanche, un satellite qui ne repasse pas souvent au même endroit servira plutôt à suivre les variations du champ de gravitation de la Terre.

Conçue pour étudier les courants océaniques, la mission franco-américaine *Topex-Poseidon* opère suivant un cycle de dix jours. Au contraire, le satellite européen *ERS-2*, spécialisé dans les mesures gravimétriques, a un cycle de 35 jours. Son prédécesseur *ERS-1* effectuait trois cycles dont les durées respectives étaient égales à 3, 35 et 168 jours. Le plus long servait surtout aux mesures géodésiques, par exemple, à la détermination des anomalies de la gravitation sous-marine qui reflètent la topographie du plancher océanique. En utilisant toutes les données altimétriques acquises par les satellites au cours des 15 dernières années, on a obtenu des cartes des fonds océaniques dont la résolution est de l'ordre du kilomètre.

Après une mission de quatre ans couronnée de succès, *ERS-1* a été « garé » en état de fonctionnement sur une orbite d'attente. Son exploitation s'est terminée le 10 mars 2000, lorsque son système de commande de position est tombé en panne. Toujours en service, son successeur *ERS-2* tourne autour de la Terre à environ 300 kilomètres d'altitude. En octobre 2001, *ENVISAT*, le plus grand satellite d'observation de la Terre jamais construit, prendra le relais. *ENVISAT* pèse huit tonnes, mesure plus de dix mètres de longueur et embarquera un altimètre de conception nouvelle et neuf capteurs de télédétection.



On déduit la distance qui sépare un satellite de la surface de la mer du temps que mettent des ondes radar pour faire l'aller-retour entre ce satellite et la surface de la mer. Pour obtenir des mesures fiables, on doit les ajuster en fonction des conditions atmosphériques, de l'humidité de la troposphère (jusqu'à 10 kilomètres d'altitude) ou encore de la teneur en électrons de l'ionosphère (de 60 à 600 kilomètres d'altitude). En outre, pour déduire des mesures de distance la hauteur des reliefs sous-marins, on doit connaître précisément la position du satellite par rapport à la Terre ou, plus exactement par rapport à un globe idéalisé, nommé ellipsoïde de référence.

létroit de Gibraltar. Des vestiges de cette dorsale médio-atlantique primitive se trouvent aujourd'hui au Sud du plateau actuel de Terre-Neuve. Puis, il y a près de 120 millions d'années, un nouveau fossé se creuse entre Terre-Neuve et la péninsule ibérique. Ensuite, le golfe de Gascogne s'ouvre, de sorte que la dorsale médio-océanique se prolonge jusqu'à la zone de fracture de Charlie Gibbs. La séparation de l'Europe de l'Ouest et de l'Amérique du Nord s'achève il y a quelque 80 millions d'années. L'avancée de la dorsale médio-atlantique se poursuit toutefois vers le Nord. Au cours de sa progression, la dorsale s'étend un moment entre le Groenland et le continent Nord-américain, qu'elle sépare, créant la mer du Labrador. Cependant, ce chemin est une impasse

et la progression de la dorsale vers le Nord s'arrête dans cette direction pour reprendre à l'Est du Groenland.

Nous avons élaboré un modèle pour décrire la période suivante. Selon cette théorie, un panache s'est élevé du manteau terrestre sous le Groenland et sous l'Europe, il y a environ 60 millions d'années. Le flux ascendant de roches chaudes s'est séparé en deux courants avant de heurter la lithosphère, à l'endroit de la zone d'expansion qui existait déjà depuis 200 millions d'années et qui avait succédé à la chaîne calédonienne, la suture de Iapetus.

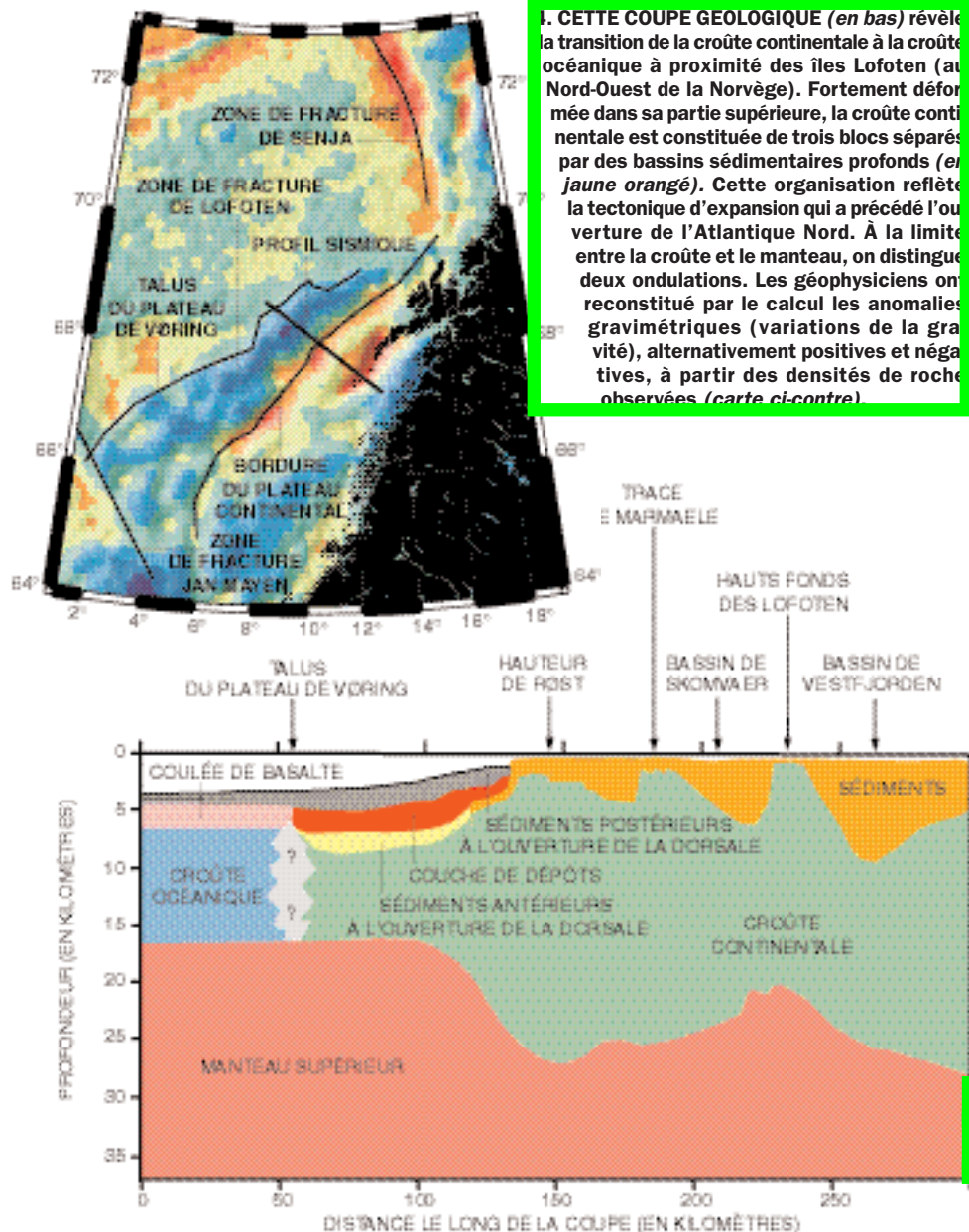
De grandes quantités de laves basaltiques (la lave est une fraction liquide de magma qui apparaît lorsque la pression ambiante diminue) remon-

èrent du manteau. Un alignement de volcans actifs apparut sur 3 000 kilomètres, du Sud de l'Angleterre au Nord de la Norvège.

L'ouverture de l'Atlantique Nord arrêta celle de la mer du Labrador qui avait commencé 36 millions d'années auparavant. Cet événement provoqua aussi probablement la transformation de la zone de fracture de Senja qui s'étend entre le Nord du Groenland et le Spitzberg en dorsale. Le Groenland se détacha alors de l'Europe et commença à se déplacer avec la plaque Nord-américaine. Le centre de la remontée des matériaux chauds issus du manteau se trouvait probablement à l'Ouest de cette nouvelle dorsale sous le plateau continental groenlandais. Cela expliquerait pourquoi cette dorsale s'est ouverte non pas au milieu de la zone d'expansion continentale de la suture de l'Iapetus, mais sur son flanc occidental : c'est au centre du panache que régnaient les contraintes les plus fortes. Sous l'effet de ces contraintes, la dorsale se déplaça vers l'Ouest dans le plateau continental groenlandais et le bloc de l'île Jan Mayen se sépara de la plaque continentale du Groenland.

Il y a environ 38 millions d'années, l'activité volcanique reprit de plus belle, notamment en Islande. Ce phénomène ne semble pas dû au panache sub-groenlandais à l'origine de l'Atlantique Nord, car les laves produites actuellement en Islande et le long de la dorsale médio-atlantique (la dorsale de Reykjanes) n'ont pas la même concentration en certains éléments chimiques rares présents à l'état de traces que les laves produites par la partie de la dorsale médio-atlantique qui s'étend plus au Nord (la dorsale de Kolbeinsey et au-delà). Du côté de l'Islande, les laves ressemblent aux volcanites de l'ancienne zone volcanique de l'Atlantique Nord, par exemple celles du plateau de Vøring, entre l'Islande et la Norvège.

Depuis près de 38 millions d'années, l'Islande est indubitablement un point chaud, c'est-à-dire le centre d'une remontée massive de magma. La dorsale médio-atlantique passe au-dessus de ce point chaud et émerge à cet endroit. Au Sud et au Nord de l'île, ce point chaud s'est déplacé de quelque 250 kilomètres vers l'Ouest au cours des derniers millions d'années, et la dorsale en a fait autant. Seule l'Islande et la zone du Vatnajokull, à l'Est de l'île



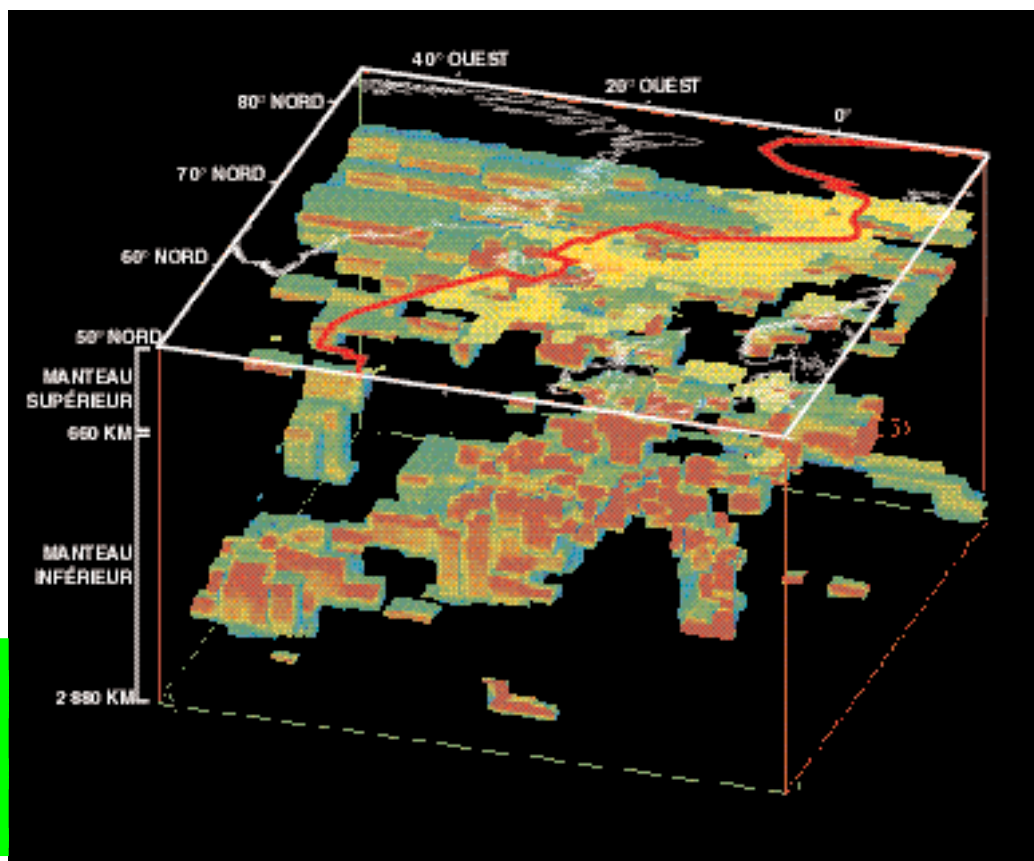
sont restées en retrait de ces mouvements d'ensemble.

L'altimétrie satellitaire nous a permis de découvrir que, lorsque l'Atlantique Nord s'est ouvert, un phénomène d'expansion a eu lieu aussi sur le plateau continental de Lofoten-Vesterålen, au large de la côte Nord de la Norvège. Nous avons été surpris de constater que les mouvements de plaques ont eu alors des conséquences pour la croûte terrestre mais aussi pour le manteau supérieur.

Dans cette région, les données topographiques indiquaient que le plancher océanique y est quasiment plat, alors que les images fournies par les satellites montraient, au contraire, des sommets et des vallées parallèles à la côte (voir la figure 4). Nous en avons conclu que les variations gravimétriques révélées par les images altimétriques ne s'expliquent pas par la forme du relief sous-marin, mais par des variations des densités des roches sous-jacentes.

Heureusement, les ingénieurs de l'industrie pétrolière avaient déjà minutieusement établi la structure géologique de la région. Trois grands bassins étroits, profonds et parallèles à la côte sont comblés, par endroits, d'une épaisseur de huit kilomètres de sédiments, de sorte que le fond de l'océan paraît quasiment plat. Ces sédiments sont plus légers que la roche de la croûte océanique (entre 5 et 20 pour cent). Les trois bassins se sont formés dans le cadre de l'expansion qui a précédé l'ouverture de l'Atlantique Nord, au début du Crétacé, il y a environ 140 millions d'années. À l'époque, la croûte s'est brisée en plusieurs blocs qui se sont affaissés. Les vallées ainsi formées se sont remplies de sédiments. Rolf Mjelde et ses collègues de l'Université de Bergen ont reconstitué deux profils sismiques du plateau de Lofoten-Vesterålen qui se chevauchent et les ont intégrés dans un modèle géologique. Poursuivant leur travail, nous avons calculé les variations gravimétriques locales en fonction des densités rocheuses. Nos résultats concordent parfaitement avec les données satellitaires.

Nos calculs et les comparaisons que nous avons faites montrent toutefois que les différences de densité dans le fond océanique ne suffisent pas à reproduire le signal gravimétrique mesuré. Pour y parvenir, il faut aussi tenir compte du sous-sol pro-



5. LA «RADIOGRAPHIE» SISMIQUE du manteau de la Terre en dessous de l'Atlantique Nord révèle des régions où les ondes sismiques se propagent à vitesse légèrement réduite. En mesurant ces vitesses, on détecte les zones où des roches particulièrement chaudes remontent des profondeurs du manteau. Un tel courant ascendant, nommé panache, a vraisemblablement joué un rôle déterminant lors de l'ouverture de la dorsale médio-atlantique (la ligne rouge). Sous cette ligne, assez près de la surface, se trouvent des zones particulièrement chaudes (en jaune). Un flux tubulaire remonte de la limite qui sépare le noyau et le manteau (au Sud-Ouest sur l'illustration) et se ramifie dans la partie supérieure du manteau, sous l'Islande. Une deuxième zone chaude, beaucoup plus large et plus plate, s'étend vers l'Ouest, sous le Groenland, jusqu'à la dorsale médio-atlantique et l'alimente probablement en magma. Le fait que les laves soient de compositions chimiques différentes au Sud et au Nord de l'Islande indique qu'il existe deux sources séparées de laves, qui correspondent peut-être aux deux structures chaudes identifiées sur l'image sismique.

fond, c'est-à-dire des ondulations de l'interface entre le manteau et la croûte terrestre. Le fait que la croûte terrestre soit si déformée à proximité du plateau continental montre la complexité des contraintes qui se sont exercées au moment de l'ouverture de la dorsale médio-atlantique, voire même avant.

Nous espérons avoir donné un aperçu de la complexité tectonique de l'Atlantique Nord. Même s'ils comprennent mieux les relations qu'entre les principales structures qui constituent le fond de l'océan Atlantique Nord, les géologues n'ont pas encore élucidé tous les mécanismes géodynamiques qui lui ont donné naissance.

Alexander BRAUN, géophysicien, travaille au Centre de recherche en géologie de Potsdam sur les données satellitaires. Gabriele MARQUART, spécialiste de géodynamique, dirige une équipe de recherche sur les interactions entre le manteau et la lithosphère, à l'Université de Francfort.

Ross TAYLOR et Scott MCLENNAN, L'évolution de la croûte continentale, in *Pour La*

Science, n° 221, mars 1996.

W. H. F. SMITH et D. T. SANDWELL, *Global Sea Floor Topography from Satellite Altimetry and Ship Sounding*, in *Science*, vol. 277, p. 1956, 1997.

Roland TROMPETTE, *Le Gondwana*, in *Pour la Science*, n° 252, octobre 1998.

Maurice MATTAUER, *Sismique et Tectonique*, in *Pour La Science*, n° 265, novembre 1999.

Les biofilms

BILL COSTERTON • PHILIP STEWART

Des bactéries organisées en films – les biofilms – résistent souvent aux antibiotiques et aux défenses immunitaires : elles causent des infections tenaces. L'étude de la structure et de l'organisation de ces communautés microbiennes aidera les microbiologistes à lutter plus efficacement contre elles.

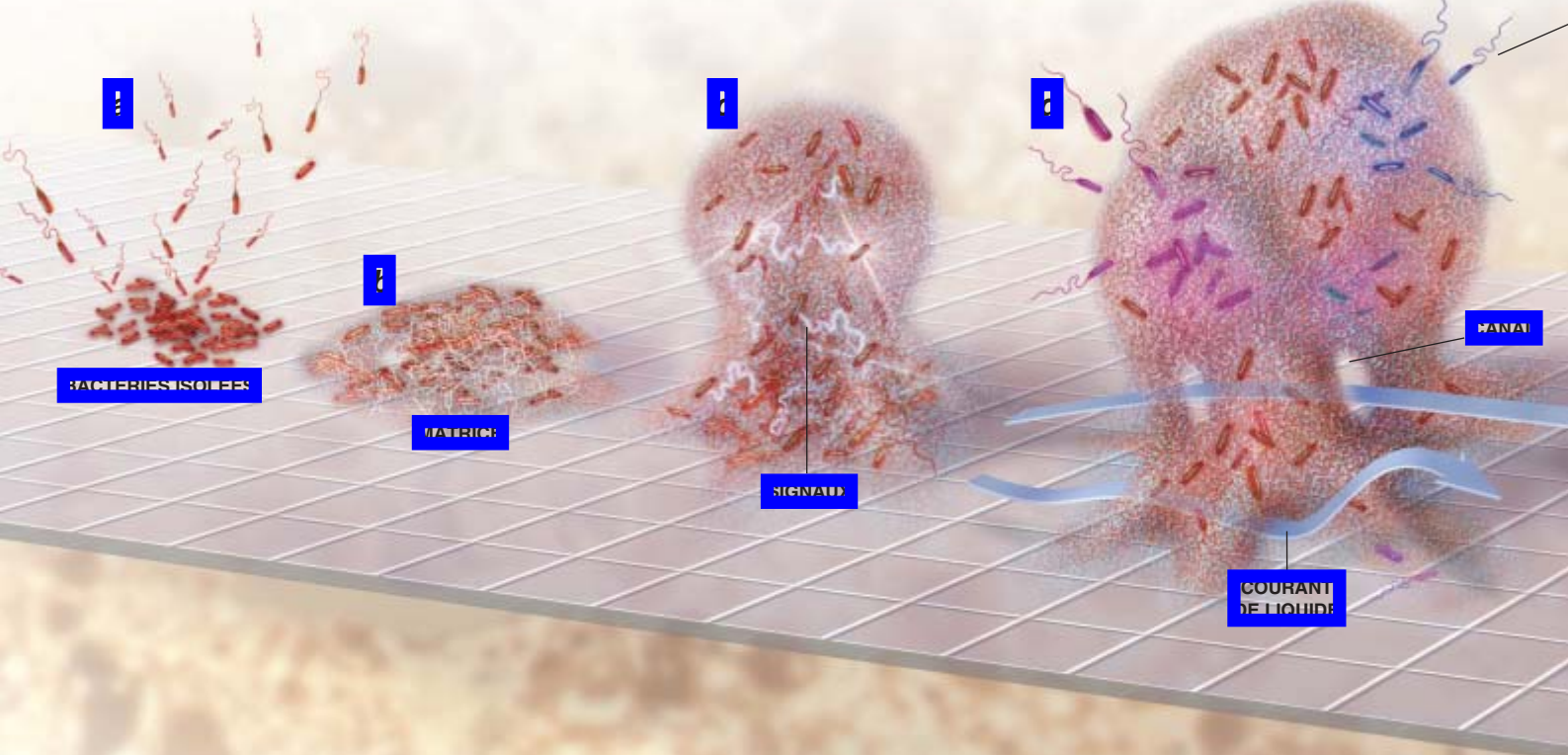
Les biofilms, des films complexes et résistants constitués de micro-organismes, sont omniprésents : dans la plaque dentaire, dans la couche glissante qui recouvre les rochers des rivières, dans l'eau des vases qui contiennent des fleurs... De plus, ils détériorent certaines installations industrielles et en contaminent les productions. Enfin, ils causent de nombreuses maladies, telles les pneumonies chez les personnes atteintes de

mucoviscidose et les infections dues aux implants artificiels. Par exemple, en 1996, une étude a montré qu'au bout d'une semaine, 10 à 50 pour cent des malades porteurs d'une sonde urinaire souffrent d'une infection due aux biofilms qui se sont développés à sa surface : au bout d'un mois, presque tous les malades sont touchés.

Pire encore, ces biofilms infectent parfois des substances thérapeutiques. En 1993 et en 1994, 100 personnes asth-

matiques sont décédées, car les aérosols qu'ils utilisaient contenaient des bactéries *Pseudomonas aeruginosa* organisées en biofilms. En 1989, les bactéries *Pseudomonas cepacia* avaient même colonisé des bouteilles d'un antiseptique puissant!

La lutte contre ce fléau est difficile, car les bactéries organisées de cette façon résistent aux antibiotiques et aux antiseptiques classiques.



1. FORMATION DES BIOFILMS : des bactéries se déplacent isolément puis se déposent sur une surface (a), s'organisent en groupes et s'y fixent. Les cellules assemblées sécrètent alors une matrice visqueuse (b), où circulent des molécules qui transmettent la consigne : se multiplier et de former une microcolonie (c). Des différences de concentrations et des courants de liquides entraînent la coexistence de plusieurs espèces bactériennes et de divers états métaboliques (d). Enfin, certaines cellules s'échappent et reprennent une vie isolée (e) peut-être pour former de nouveaux biofilms. Chacune de ces étapes est une cible possible pour lutter contre la formation des biofilms.

Une des clefs de cette redoutable caractéristique tient aux systèmes de communication utilisés par les micro-organismes des biofilms. Les bactériologues mettent peu à peu au jour les mécanismes moléculaires de ces biofilms, et comprennent comment, de l'organisation elle-même, naît la résistance des micro-organismes. Ces mécanismes élucidés, ils pourront concevoir des substances, à usage thérapeutique ou industriel, qui interféreront avec les signaux de transmission et limiteront les dégâts dus à ces biofilms.

Des bactéries organisées

Pourquoi les biofilms ont-ils si longtemps gardé leurs secrets? En partie parce que les méthodes d'observations n'étaient pas assez puissantes. Depuis la fin du XIX^e siècle et les travaux de Robert Koch qui confirmaient l'origine microbienne de diverses maladies, les bactéries sont considérées comme des cellules uniques qui flottent ou nagent dans un liquide. Les

biologistes ont d'abord observé des cellules de culture mises en suspension dans une goutte de liquide à l'aide de microscopes optiques. Toutefois, ces conditions expérimentales ne reflètent pas l'environnement réel, et le comportement des bactéries cultivées en laboratoire n'a rien de commun avec celui des bactéries dans la nature.

Le microscope électronique n'a pas été d'un grand secours, car ceux qui l'ont utilisé n'ont pas obtenu d'images précises des couches profondes des amas de bactéries : ils en avaient conclu que les cellules situées à l'intérieur étaient mortes et disposées n'importe comment.

Ce point de vue a changé dans les années 1990, lorsque les bactériologues ont utilisé la microscopie confocale à balayage laser : on fait un balayage de l'échantillon au moyen d'un faisceau laser dont on peut ajuster la profondeur. On obtient ainsi une

série de « coupes » optiques transversales dont on peut reconstituer la profon-

deur varie et l'on reconstruit ensuite une image en trois dimensions. Depuis on a découvert que ces micro-organismes adhèrent aux surfaces humides sous forme de colonies organisées aux caractères variés.

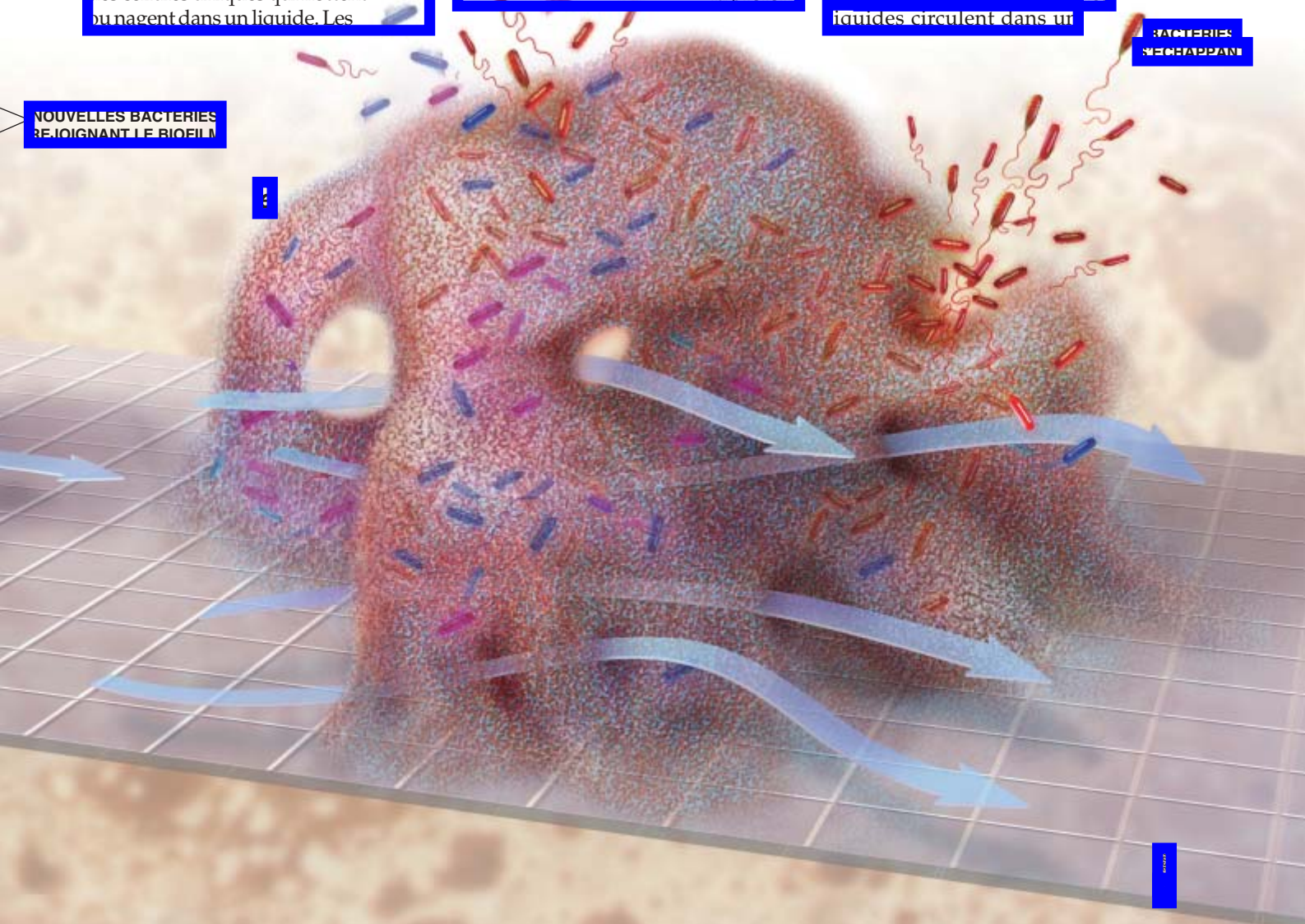
Notons que les bactéries ne sont pas les seules à créer des biofilms (voir la figure 4) : les champignons, les algues, les protozoaires s'organisent aussi en biofilms. Cette diversité relève sans doute d'une stratégie ancienne de croissance.

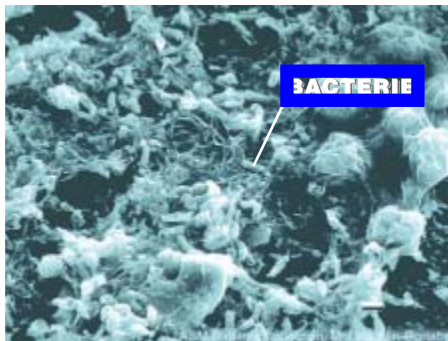
L'étude des biofilms a montré que les bactéries poussent dans de minuscules enclaves, des microcolonies, qui en grand nombre, constituent les biofilms. Les bactéries elles-mêmes représentent moins d'un tiers de l'ensemble : le reste est une substance visqueuse, nommée matrice, sécrétée par les cellules, qui assure la cohésion de chaque microcolonie, la protège, absorbe l'eau et piège les petites particules circulantes.

Au sein d'un biofilm, les liquides circulent dans un

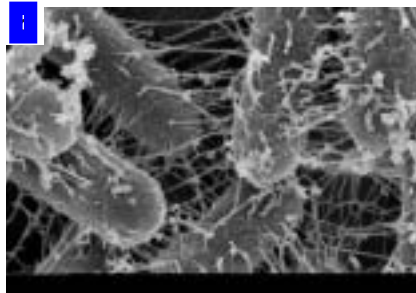
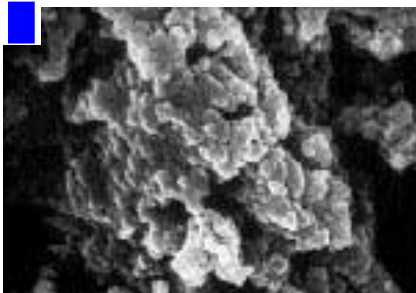
BACTERIES
SECHAPPAIN

NOUVELLES BACTERIES
REJOIGNANT LE BIOFILM





2. LES LENTILLES DE CONTACT sont des surfaces aisément colonisées par des biofilms, c'est-à-dire par des micro-organismes. (la photo en médaillon) entourés d'une matrice organique.



3. DES BACTERIES *PSEUDOMONAS AERUGINOSA* cultivées dans un tube en Téflon sont observées au microscope électronique après un jour (a), trois jours (b) et six jours (c). Le biofilm qu'elles fabriquent ressemble à ceux que ces bactéries forment dans les voies respiratoires lors d'infections chroniques (d).

réseau de minuscules conduits et baignent chaque microcolonie de bactéries. Ils fournissent ainsi des nutriments dissous et éliminent leurs déchets comme le liquide intérieur d'un organisme pluricellulaire.

En raison de cette organisation, les cellules de la périphérie d'une microcolonie sont mieux desservies par le réseau de canaux que celles du centre. La densité des cellules et la présence de la matrice organique s'opposent à l'écoulement de l'eau autour des cellules du centre qui reçoivent seulement les nutriments qui diffusent jusqu'à elles. Bien que les petites molécules circulent facilement dans la matrice – essentiellement constituée d'eau –, une substance ne diffuse jusqu'au centre d'une microcolonie que si, sur son parcours, elle interagit peu avec les cellules et avec la matrice.

Ainsi, dans un biofilm, l'environnement des cellules varie notablement. En 1985, grâce à des micro-électrodes, on a constaté que la concentration en oxygène diffère entre deux positions situées à une distance de cinq centièmes de millimètre, soit à peine plus que le diamètre d'un cheveu ! La quantité d'oxygène d'une communauté bactérienne reflète la vitalité du métabolisme et, par conséquent, l'état physiologique des cellules. Par exemple, dans un biofilm constitué de *Pseudomonas aeruginosa* (la bactérie responsable des pneumonies associées à la mucoviscidose), les bactéries sont actives et se multiplient seulement là où l'oxygène pénètre, c'est-à-dire sur deux ou trois centièmes de millimètre, en périphérie de chaque minuscule colonie. Plus en profondeur, les cellules sont en dormance : elles sont vivantes, mais en repos (voir la figure 5). Contrairement à l'uniformité qui règne dans les cultures de laboratoire, la diversité des états métaboliques prévaut dans les conditions naturelles.

A l'intérieur d'un même biofilm, une cellule peut se comporter différemment d'une cellule voisine génétiquement identique. Par ailleurs, les conditions locales influent sur le pouvoir pathogène des micro-organismes d'un biofilm : certains font peu de dégâts chez un hôte, d'autres sont mortels. La variété des conditions rend également possible la cohabitation d'espèces différentes. Certaines, par exemple, se nourrissent des déchets des autres.

La digestion des ruminants, telles les vaches, est un exemple d'une telle coopération. Le fourrage que les ruminants

ingèrent passe dans les deux premières cavités de l'estomac (la panse et le rumen), où se développent des biofilms constitués de bactéries qui digèrent la cellulose (un polymère de glucose) des plantes et produisent des acides gras. Quand ces derniers sont trop nombreux, ils inhibent la croissance des bactéries qui les ont fabriqués. Des cellules mobiles de *Treponema* (des bactéries hélicoïdales et allongées) et d'autres espèces envahissent alors le biofilm et en consomment les constituants. Les matériaux du fourrage sont peu à peu convertis en une masse de bactéries que l'animal digère, après régurgitation, dans deux autres cavités, le feuillet et la caillette. Ainsi, les vaches se nourrissent de biofilms bactériens et non de foin!

Pour les ruminants, les biofilms sont indispensables. Toutefois, ils constituent parfois une nuisance, voire un réel danger pour la santé des êtres humains, car ils résistent à la plupart des traitements bactéricides utilisés en médecine et dans l'industrie, alors que ces mêmes traitements détruisent les bactéries lorsqu'elles sont isolées. Les bactéries des biofilms échappent aussi aux molécules et aux cellules du système immunitaire.

La cohésion fait la force

Pourquoi les biofilms sont-ils aussi résistants (voir la figure 5)? D'abord parce que les substances antimicrobiennes doivent diffuser dans une matrice compacte : par exemple, la concentration d'un antibiotique est plus faible au cœur de la matrice qu'à l'extérieur. On doit donc utiliser des quantités importantes de substances bactéricides, au risque de dépasser la dose tolérée. Par exemple, les antibiotiques de la famille des pénicillines pénètrent difficilement dans les biofilms constitués de staphylocoques, car ces bactéries fabriquent une enzyme nommée bêta-lactamase, qui dégrade l'antibiotique plus rapidement que celui-ci ne diffuse dans le biofilm. L'antibiotique n'atteint jamais les couches les plus profondes.

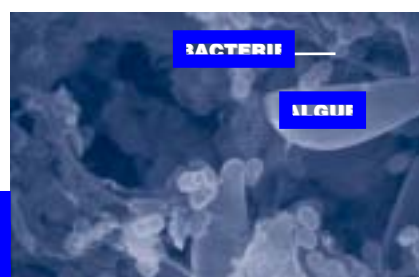
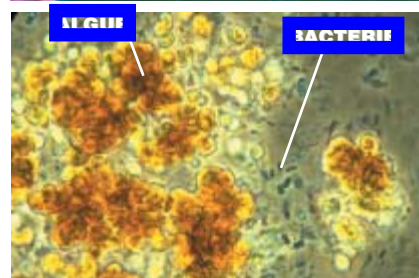
Les produits chlorés, prisés pour des usages domestiques ou industriels, sont également peu efficaces. Cet oxydant doit lutter, couche par couche, pour neutraliser le biofilm. Ce processus est long et nécessite beaucoup de produit chloré.

De surcroît, bien que certains agents antimicrobiens pénètrent assez bien

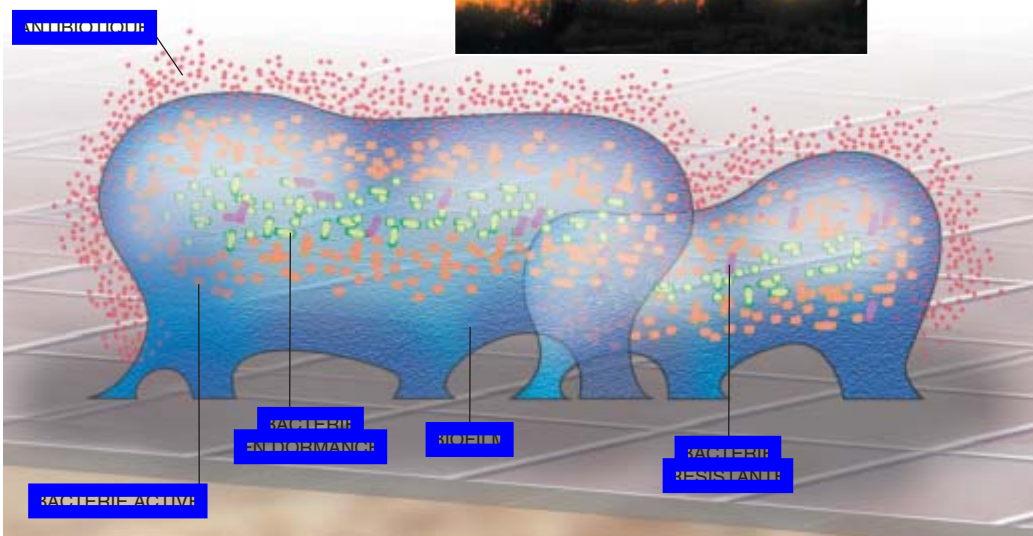
dans les biofilms, les micro-organismes en amas survivent souvent aux traitements agressifs qui auraient éradiqué des cellules isolées. Les biologistes ont récemment découvert comment les conditions d'organisation des bactéries dans un biofilm les protègent contre les agents antibactériens.

La pénicilline agit au moment où une bactérie se divise en deux cellules filles, lorsque chacune construit une paroi protectrice, nommée peptidoglycane, plus ou moins épaisse selon les types de bactéries, et constituée de dérivés de sucre et d'acides aminés reliés entre eux. L'antibiotique empêche certaines liaisons de s'établir et déstabilise l'ensemble de la structure. Dépourvues de leur paroi, les cellules bactériennes sont rapidement détruites. Or, au cœur d'un biofilm, des bactéries sont parfois en état de dormance : elles sont vivantes, mais ne se divisent pas, et, par conséquent, sont insensibles à la pénicilline. Dans la mesure où des micro-organismes actifs et inactifs coexistent dans un biofilm et où les bactéries survivantes peuvent se nourrir des bactéries mortes, les quelques cellules qui ont résisté à une antibiothérapie sont capables de reconstituer un biofilm en quelques heures.

Ainsi, des antimicrobiens efficaces sur des cellules en culture sont sans effet sur les biofilms qui causent des maladies ou qui, dans les installations



4. LES BIOFILMS sont parfois constitués de protozoaires, des organismes unicellulaires (en haut), ou bien d'algues (au centre, et en bas ici elles sont avec des bactéries).



5. L'ORGANISATION D'UN BIOFILM confère aux bactéries une résistance accrue aux substances bactéricides pour plusieurs raisons. D'abord, la matrice réduit la concentration en antibiotiques (en rouge) autour des bactéries. Par ailleurs, certaines bactéries au centre du biofilm sont à l'état de dormance (en vert) et sont donc moins sensibles aux antibiotiques. Cet état de dormance des bactéries peut être mis en évidence par un colorant : sur la photographie, une croissance rapide apparaît en orange et une croissance ralentie, en vert. Enfin, protégées, mais exposées aux substances bactéricides, les bactéries peuvent développer une résistance transmissible (en rose).

industrielles, ont des effets désastreux. Ils encrassent les machines et accélèrent la corrosion des canalisations en métal (voir la figure 6).

Bactéries communicantes

On a récemment mis en évidence que lorsque des bactéries adhèrent à une surface et forment un biofilm, elles fabriquent des centaines de protéines que ne synthétisent pas les bactéries isolées : certaines de ces protéines participent au glissement des cellules sur une surface avant qu'elles ne s'immobilisent. De plus, chez *Staphylococcus epidermidis*, la bactérie responsable d'infections à staphylocoques, les gènes qui commandent l'étape suivante du développement d'un biofilm, c'est-à-dire la synthèse de la matrice extracellulaire, sont inactivés. Lorsque ces gènes sont inactivés, les bactéries deviennent

incapables de former des biofilms dans un tube à essai ou dans les tissus d'animaux de laboratoire.

Des expériences récentes ont révélé l'existence de gènes similaires chez d'autres espèces. Ainsi, chez *Pseudomonas aeruginosa*, plusieurs gènes sont inactivés dans les 15 minutes qui suivent l'attachement de la bactérie à une surface. L'un de ces gènes, *algC*, est nécessaire à la synthèse d'alginate, le polymère gélatineux qui constitue une grande partie de la matrice extracellulaire. L'alginate, aussi fabriqué par les algues, est utilisé dans l'industrie alimentaire.

Comment les bactéries qui se rassemblent pour former un biofilm « décident »-elles quels gènes activer ? Ces micro-organismes, en apparence autonomes, communiquent entre eux. Chez *Pseudomonas aeruginosa* et dans une vaste classe de bactéries analogues, les molécules de

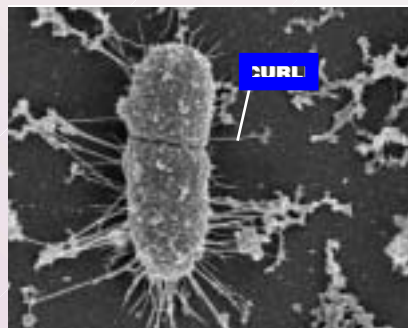
signalisation sont des homosérines lactones acylées. Ces molécules, fabriquées par chaque cellule en petite quantité, diffusent librement à travers les enveloppes cellulaires. Leur concentration augmente avec la densité cellulaire : au-dessus d'un seuil, les gènes à l'origine de la matrice sont activés, et le biofilm apparaît. Par ailleurs, ces homosérines lactones acylées sont aussi à l'origine de la synthèse de facteurs de toxicité, tel l'acide cyanhydrique. Les bactéries *Pseudomonas aeruginosa* dépourvues du gène qui code les enzymes fabriquant ces lactones n'édifient pas de biofilms et s'amoncellent en amas désorganisés.

Les biologistes connaissent ainsi les molécules de signalisation des biofilms qui se développent, notamment à l'intérieur des sondes urinaires. Ces biofilms ainsi que ceux qui s'installent sur les implants permanents

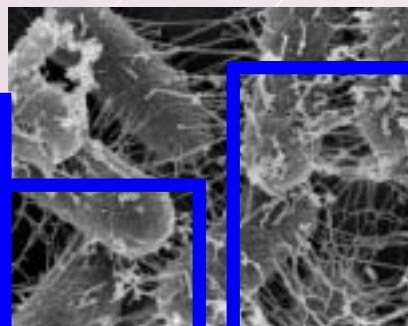
La perception des bactéries

De nombreux microbiologistes luttent contre les biofilms en recherchant des molécules qui interféreraient avec les mécanismes de communication nécessaires au bon développement de ces structures pluricellulaires. De telles molécules empêcheraient le biofilm d'acquiescer l'organisation responsable de ses propriétés de résistance. Les défenses immunitaires ou des traitements antimicrobiens viendraient alors aisément à bout de ces structures « inachevées ». On cherche aussi à mettre au point des revêtements de surface qui évitent la contamination des biomatériaux implantables ou des installations industrielles. Cette stratégie a été élaborée après que l'on a compris comment les bactéries qui nagent dans un liquide « savent » qu'elles ont heurté une surface solide.

Avec mon équipe, nous avons constaté que la synthèse des appendices grâce auxquels certaines bactéries *Escherichia coli* s'attachent aux matériaux est amplifiée après quelques instants de contact (voir la figure). Ces appendices sont de petits poils, nommés curli. C'est aussi le cas de la sécrétion de l'acide colanique, un polysaccharide qui confère sa consistance à la matrice, comme les pectines des confitures. À l'inverse, les appendices, nommés flagelles, que la bactérie utilise pour nager dans la phase liquide disparaissent progressivement à mesure que les bactéries se divisent, car ils ne sont pas renouvelés. Cette transition physiologique résulte de la perception à l'interface solide-liquide d'un environnement différent de la phase liquide. Chez les bactéries Gram négatives (les bactéries sont différenciées selon leur couleur après une coloration de Gram : elles sont Gram positif ou négatif), nous avons identifié des capteurs des protéines qui traversent la membrane et



CURLI



LES BACTÉRIES *ESCHERICHIA COLI* qui s'organisent en biofilms modifient leurs appendices : quand elles se fixent sur une surface, les curli sont rares, mais ils deviennent plus nombreux à mesure que le biofilm croît.

qui participent à cette perception. La conformation du domaine protéique en contact avec le milieu extérieur dépend des paramètres ioniques et électriques de ce milieu (acidité, différence de concentration de molécules entre l'intérieur et l'extérieur, présence de charges électriques). Au voisinage d'une surface, les interactions avec les molécules d'eau et les ions en solution modifient ces paramètres. Les changements extérieurs de conformation de la protéine entraînent des changements de la partie intracellulaire du capteur, puis l'information parvient au génome par l'intermédiaire d'autres protéines régulatrices qui sont phosphorylées.

En collaboration avec Gilbert Legeay, du Centre de transfert de technologie du Mans, nous avons utilisé ces résultats pour mettre au point des revêtements antibiofilms destinés aux biomatériaux implantables. Quand on greffe des molécules chargées – inoffensives pour l'organisme humain – à la surface de ces matériaux, les bactéries ne perçoivent plus l'interface de la même façon et ne se fixent pas immédiatement, mais au bout de quelques dizaines d'heures. Dans le cas d'implants posés pour des périodes courtes, tels les sondes urinaires ou les cathéters intravasculaires, ce répertoire évite de nombreuses contaminations. Lorsque le risque de complications infectieuses est lié aux germes introduits accidentellement au moment de l'opération, par exemple, lors de l'implantation de prothèses ou de cristallins artificiels, les défenses immunitaires ou une antibiothérapie appropriée peuvent lutter contre des bactéries isolées et éviter l'infection avant que les implants n'aient été colonisés par des biofilms plus résistants.

Philippe LEJEUNE, CNRS-UMR 5571
Institut national des sciences appliquées de Lyon

causent des infections difficiles à traiter. Ces infections insidieuses, lentes à se développer, se révèlent parfois par des crises soudaines et répétées et sont difficiles à éradiquer. Par ailleurs, les biofilms ont également été incriminés dans certaines infections parodontales, des infections de la prostate, des calculs rénaux, la tuberculose, la maladie du légionnaire et certaines infections de l'oreille moyenne.

Pour éviter la formation de ces biofilms, on pourrait, par exemple, utiliser des molécules qui, en recouvrant la surface des cellules, neutraliseraient leurs propriétés adhésives, ce qui réduirait leur aptitude à se fixer (voir l'encadré page 52). On pourrait aussi bloquer la synthèse de la matrice extracellulaire ou encore inhiber la synthèse des molécules que les bactéries des biofilms utilisent pour communiquer ; on arrêterait ainsi la formation des biofilms, on éviterait la production de toxines et l'on bloquerait diverses activités néfastes.

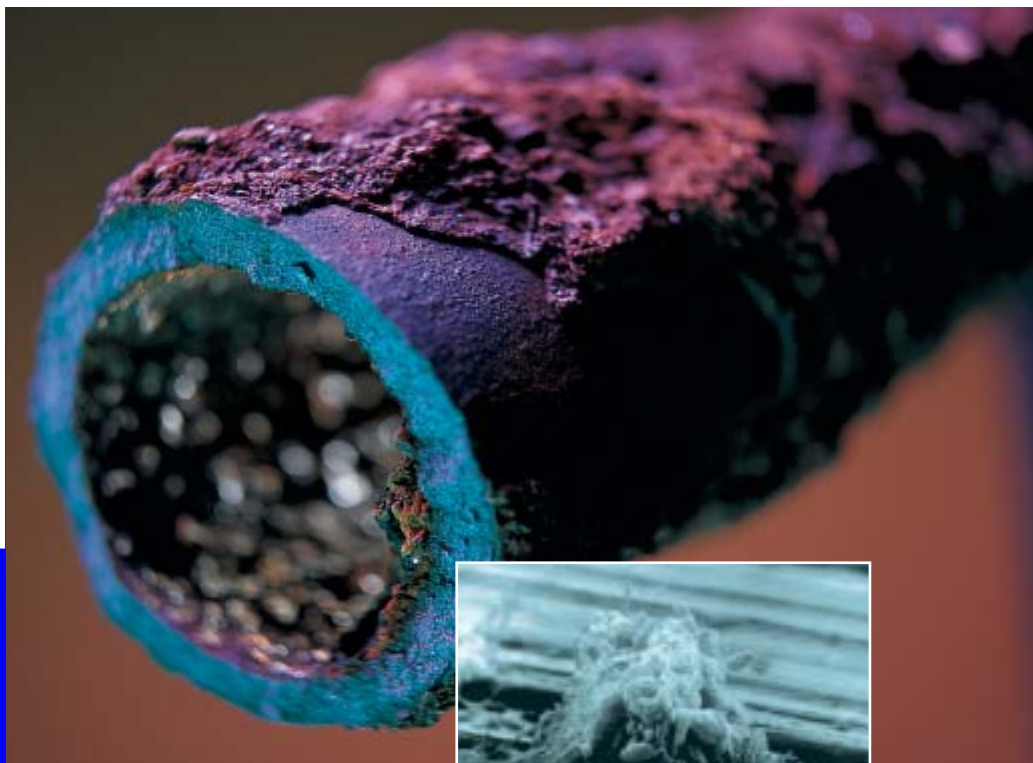
Plutôt que de submerger les organismes agressifs avec des poisons, mieux vaut sans doute manipuler les cellules de façon plus subtile. En inhibant diverses activités spécifiques de ces micro-organismes, on lutte plus efficacement qu'en utilisant des doses massives de substances, toxiques pour les bactéries, mais aussi pour d'autres qui sont inoffensives, voire utiles au fonctionnement de l'organisme.

Le salut par les algues

En 1995, Staffan Kjellerberg et Peter Steinberg, de l'Université de Sydney, ont remarqué que les feuilles d'une algue rouge (*Delisea pulchra*) sont rarement recouvertes de biofilms. Malgré les milliers d'espèces bactériennes qui vivent dans l'eau de mer, ces algues restent immaculées. Elles utilisent des substances chimiques de la famille des furanones.

Les furanones ressemblent aux homosérines lactones acylées et à des molécules, récemment découvertes, qui permettent aux bactéries et à des espèces différentes de communiquer. Les furanones semblent se lier aux récepteurs des bactéries qui, normalement, fixent les signaux déclenchant la fabrication des biofilms : en empêchant la fixation de ces signaux, les furanones évitent la formation des biofilms.

Ces composés inhibent la formation de nouveaux biofilms, mais détruisent aussi les biofilms existants. Les furanones



5. LES BIOFILMS DIMINUENT LE RENDEMENT
des installations industrielles, voire les endommagent : par exemple, ils peuvent accélérer la corrosion des canalisations métalliques.

ones sont adaptées à un usage médical puisqu'elles sont dépourvues de toxicité et assez stables dans l'organisme. Par ailleurs, bien que les furanones existent dans les océans depuis des millions d'années, aucune souche bactérienne résistante ne semble avoir été sélectionnée. On peut espérer que les bactéries colonisant les dispositifs médicaux et les tissus humains ne deviendront pas non plus résistantes.

Aujourd'hui, il existe des revêtements protecteurs qui contiennent des furanones et qui sont appliqués sur les coques de bateaux et les équipements d'aquaculture.

Selon certains biologistes, la formation des biofilms bactériens s'apparente à un processus d'embryogenèse.

de même qu'un œuf fécondé donne naissance à divers types de cellules au cours du développement d'un fœtus, de même certaines bactéries se différencient après leur ancrage à une surface, synthétisent des molécules de communication (qui rappellent les phéromones et les hormones) et coordonnent l'édification de microcolonies dont la structure est complexe. Par ailleurs, dans certains biofilms, des bactéries d'espèces différentes cohabitent, certaines digèrent, par exemple, des nutriments pour une autre espèce. Les bactéries, longtemps considérées comme des micro-organismes isolés et rudimentaires, semblent dotées de capacités d'organisation et de coopération notablement plus élaborées.

Bill COSTERTON est directeur du Centre d'études des biofilms à l'Université du Montana.

Philip STEWART est bactériologiste au Centre d'études des biofilms.

J. W. COSTERTON, P. STEWART et E. F. GREENBERG, *Bacterial Biofilms : A Common Cause of Persistent Infections*, in *Science*, vol. 284, pp. 1318-1322, mai 1999.

D. G. ALLISON, P. GILBERT, H.M. LAPIN-SCOTT et M. WILSON, *Community*

Structure and Co-operation in Biofilms, Cambridge University Press, 2001.

B. COSTERTON, *Cystic Fibrosis Pathogenesis and the Role of Biofilms in Persistent Infection*, *Trends in Microbiology*, vol. 9, n° 2, pp. 50-52, 2000.

D. NIVENS, D. OHMAN, J. WILLIAMS et M. FRANKLIN, *Role of Alginate and Alginate O-Acetylation in the Formation of Pseudomonas Aeruginosa Microcolonies and Biofilms*, n° 183, pp. 1047-1057, 2001.

L'Univers de Georges Lemaître

DOMINIQUE LAMBERT

Prêtre et physicien, Georges Lemaître fut l'un des fondateurs de la théorie du Big Bang. Nombre de ses intuitions, qu'ils défendit parfois contre Einstein lui-même, se sont révélées, près de 50 ans plus tard, d'une portée fondamentale.

En 1933, Albert Einstein donna une série de cours à la Fondation universitaire de Bruxelles. Lorsqu'un collègue lui demanda si tous les scientifiques l'avaient bien compris, il répondit : «Le professeur De Donder peut-être, le chanoine Lemaître sûrement. Les autres je ne crois pas.»

Georges Lemaître est considéré comme l'un des fondateurs de la théorie du Big Bang sur laquelle repose la cosmologie moderne. Parce qu'il fut aussi un homme de foi, certains ont prétendu que l'hypothèse d'un cataclysme créant l'Univers à une date déterminée dans le passé constituait de sa part une justification scientifique de la «création» du monde selon les chrétiens.

Ainsi, l'astronome britannique Fred Hoyle, partisan, pour des raisons philosophiques, d'un modèle d'Univers éternel, forgea-t-il le terme péjoratif de Big Bang («gros boum» en anglais) pour ridiculiser les idées développées et popularisées par Lemaître. Ironie du sort, ce bon mot est aujourd'hui utilisé sans connotation négative pour désigner une théorie que de nombreux faits expérimentaux ont, depuis lors, étayée. Quant aux convictions scientifiques de Lemaître, elles se fondaient non pas sur sa foi (il sut toujours éviter toute confusion entre sciences et religion) mais sur des arguments mathématiques et physiques de grande valeur. Quelques moments marquants de sa carrière éclairent la portée scientifique de certaines de ses intuitions.

Un prêtre à Cambridge

Georges Lemaître naît le 17 juillet 1894 à Charleroi, en Belgique. En 1911, il commence des études d'ingénieur des mines pour lesquelles il montre peu d'enthousiasme et, à son retour de la Première Guerre mondiale, il change d'orientation pour suivre un enseignement de mathématiques et de physique qu'il termine en 1920. La même année, il entre au séminaire de Malines où, tout en se préparant à devenir prêtre, il continue à étudier les ouvrages traitant de relativité restreinte et générale. Il rédige un mémoire intitulé *La physique d'Einstein* qui lui vaut de gagner un concours de bourse d'études et lui permet, après son ordination le 22 septembre 1923, de partir étudier une année à Cambridge, en Grande-Bretagne. Là, Lemaître travaille sous la direction d'Arthur Eddington, l'astronome qui cinq ans auparavant, a confirmé ce qu'Einstein avait prévu : des rayons lumineux passant au voisinage du Soleil sont déviés par la force de gravitation.

En 1924-1925, il parachève ses études à l'Institut de technologie du Massachusetts et visite certains des hauts lieux de l'astronomie mondiale, notamment l'Observatoire du mont Palomar, où se trouve alors le plus grand télescope jamais construit. Ainsi, le jeune expert des nouvelles théories de l'espace-temps entre en contact avec l'astronomie au moment même où la cosmologie scientifique va naître.

Dans quel «contexte cosmologique» Lemaître se trouve-t-il ? Au milieu des années 1920, les astronomes attribuent à l'Univers observé une taille de quelques dizaines de milliers d'années-lumière, c'est-à-dire six ordres de grandeur inférieure à celle qu'on lui reconnaît aujourd'hui (de l'ordre de dix milliards d'années-lumière). Par ailleurs, en vertu d'un préjugé hérité du XIX^e siècle, on n'envisage pas que cet Univers évolue, encore moins qu'il ait un âge fini.

Quant à son contenu, depuis le XVIII^e siècle, les astronomes pensent que les étoiles visibles dans le ciel se rassemblent en un vaste disque plat, la Voie lactée, au sein duquel se trouve aussi le Soleil. Dès 1785, le philosophe allemand Emmanuel Kant avait posé l'hypothèse que les nébuleuses spirales découvertes par les astronomes à l'aide de leurs premiers télescopes étaient de gigantesques regroupements d'étoiles semblables à la Voie lactée. Aujourd'hui nous nommons galaxies ces «univers» imaginés par le philosophe. Toutefois, jusqu'au milieu des années 1920, cette hypothèse est considérée comme douteuse ; en effet, elle suppose que ces nébuleuses doivent se trouver à des millions d'années-lumière. Or, les astronomes y observent régulièrement des explosions lumineuses qui, si elles étaient situées à de telles distances, devraient dégager en très peu de temps des quantités d'énergie qu'aucun des mécanismes physiques connus à l'époque n'est capable de produire.

Aujourd'hui, on sait qu'il s'agit de supernovae, des explosions d'étoiles géantes au cours desquelles des réactions thermonucléaires fournissent en quelques heures plus de lumière que des milliards d'étoiles.

Un Univers éternel

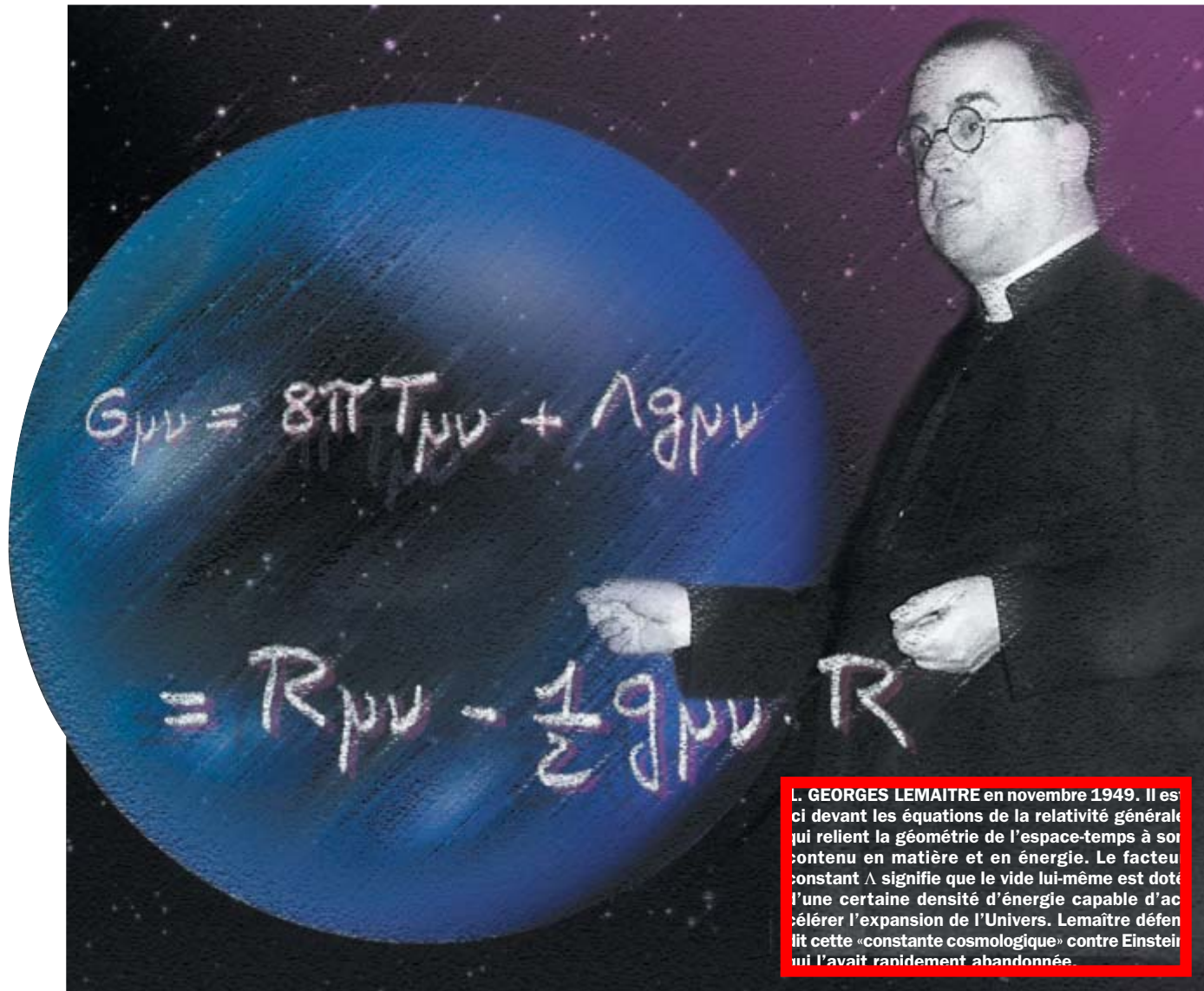
L'Univers des astronomes de l'époque aussi loin qu'ils l'observent, ne contient qu'une seule galaxie, la nôtre. Est-il clos et de dimension finie, ou bien infini et, par conséquent, presque vide? À partir de 1916, cette question prend tout son sens avec la publication par Albert Einstein de sa théorie de la relativité générale, qui permet aux astronomes de construire différents modèles cosmologiques selon les hypothèses choisies. Einstein a montré que la force de gravitation peut être considérée comme une manifeste

ation de la courbure de l'espace-temps. Ses équations permettent de calculer cette courbure en chaque point de l'Univers, connaissant la quantité de matière (ou d'énergie) qu'il y trouve. Quant à la géométrie globale de l'Univers, il est possible, moyennant quelques hypothèses sur la répartition de la matière et de l'énergie en son sein, d'en donner une description rigoureuse.

Le premier modèle d'Univers proposé par Einstein est une solution de ses équations où le cosmos est clos et statique (voir la figure 3). Dans ce modèle, l'espace est une «hypersphère» de rayon constant (voir la figure 2). De même qu'une créature plate vivante «dans» la surface (à deux dimensions) d'une sphère verrait son Univers comme un espace courbe, fini mais sans frontières, nous vivrions dans un espace à trois dimensions qui serait la

«surface» d'une hypersphère (une sphère dans un espace à quatre dimensions). Pour que le rayon de cet Univers reste constant (c'est-à-dire pour empêcher que le poids de la matière ne provoque son effondrement), Einstein a postulé l'existence d'une «force répulsive» capable de contrecarrer les effets de la gravitation et de maintenir l'Univers en équilibre. Cette force intervient dans les équations de la relativité générale sous la forme d'une constante, connue sous le nom de «constante cosmologique».

Peu de temps après, l'astronome hollandais Willem De Sitter, trouve un autre modèle d'Univers, également solution des équations d'Einstein, où l'espace est infini et la densité de matière nulle. Lorsque Lemaître commence à travailler sur les problèmes cosmologiques, la communauté des astronomes ne se réfère qu'à ces deux modèles.



L. GEORGES LEMAÎTRE en novembre 1949. Il est ici devant les équations de la relativité générale qui relient la géométrie de l'espace-temps à son contenu en matière et en énergie. Le facteur constant Λ signifie que le vide lui-même est doté d'une certaine densité d'énergie capable d'accélérer l'expansion de l'Univers. Lemaître défendit cette «constante cosmologique» contre Einstein qui l'avait rapidement abandonnée.

Un Univers en expansion

Cependant, en 1925, à l'Observatoire du mont Palomar, l'astronome américain Edwin Hubble montre que la «nébuleuse» d'Andromède se situe, en fait, à plusieurs millions d'années-lumière de nous ; la polémique sur la nature des «nébuleuses» est réouverte. Georges Lemaître prend donc contact avec la cosmologie au moment même où ressuscite l'idée d'un Univers immense et parsemé de galaxies. En outre, les astronomes sont en train d'établir que la lumière que nous recevons de la plupart de ces galaxies est décalée vers les grandes longueurs d'ondes (vers le rouge), ce qui semble indiquer qu'elles s'éloignent de nous rapidement.

Ainsi, Lemaître est naturellement conduit à chercher une explication à cette «récession» des galaxies (indépendamment du russe Alexander Friedman). En 1927, il propose une troisième solution des équations d'Einstein : la taille de l'Univers croîtrait de façon exponentielle et se raccorderait, aux deux grands modèles de l'époque, dans un passé et dans un futur infiniment éloignés (voir la figure 4). Aux temps très reculés, ce Univers se comporterait comme l'Uni-

vers statique d'Einstein. Vers le futur, l'Univers de Lemaître, dont la masse est constante et dont la taille grandit indéfiniment, tendrait vers le modèle de De Sitter, totalement vide. Entre ces deux extrêmes, l'expansion de l'Univers expliquerait que les galaxies se fuient les unes les autres.

Au moyen de cette solution, dès 1927, Lemaître établit la loi qui gouverne la récession des galaxies ; elle stipule que deux galaxies se fuient à une vitesse proportionnelle à la distance qui

2. UNE CREATURE A DEUX DIMENSIONS, vivant dans la surface de cette sphère, constate que son univers est fini (elle peut en faire le tour), mais sans frontières. La géométrie n'y est pas euclidienne : la somme des angles d'un triangle (en bleu) est supérieure à 180 degrés, et l'on y trouve des «droites» (le plus court chemin entre deux points) qui sont parallèles mais se rejoignent à une distance finie (par exemple deux perpendiculaires à l'équateur de la sphère se croisent au pôle). Toutefois, à petite échelle, la géométrie est très proche de la géométrie euclidienne. Beaucoup de modèles d'Univers stipulent que l'espace, euclidien à notre échelle, est la surface d'une hypersphère, un analogue à trois dimensions du monde de la créature plate.

les sépare. A partir d'un catalogue de 42 galaxies dont on connaît un ordre de grandeur des distances ainsi que les vitesses de fuite, il estime la constante de proportionnalité à 625 kilomètres par seconde et par mégaparsec (c'est-à-dire que deux galaxies distantes de un mégaparsec, soit un peu plus de trois millions d'années-lumière, se fuient à 625 kilomètres par seconde). Ainsi, Lemaître est le premier à établir théoriquement la loi... de Hubble même si la valeur qu'il donne à la constante, dite elle aussi «de Hubble» est surévaluée. L'astronome américain eut les honneurs de la postérité, car il fut le premier, deux ans plus tard, à publier une compilation détaillée d'observations montrant cette loi.

Avec le modèle de 1927, Lemaître est l'un des tous premiers cosmologistes à envisager que l'Univers évolue. Toutefois, l'Univers n'y a pas encore de commencement. Pourquoi Lemaître privilégie-t-il – momentanément – l'idée d'un Univers au passé infini ? Il semble que ce choix découle de sa surestimation de la valeur de la constante de Hubble. En effet, si l'Univers est aujourd'hui en expansion très rapide, sa taille, dans un passé relativement récent, devait être beaucoup plus petite. En fait, l'inverse de la constante de Hubble donne un ordre de grandeur de la période récente d'expansion. Avec une valeur de l'ordre de 625 kilomètres par seconde et par mégaparsec, Lemaître trouve que l'Univers a moins d'un milliard d'années, ce qui est inférieur aux deux milliards d'années que l'on attribue à la Terre à son époque. Cette évaluation

3. UN MODELE D'UNIVERS ETERNEL avait la préférence d'Einstein. Dans ce cadre, la taille – finie – de l'Univers est constante. Comme cet Univers a tendance à s'effondrer sur lui-même parce que les masses s'attirent sous l'effet de la gravitation, Einstein postula l'existence d'une seconde force, répulsive, qu'il introduit dans ses équations sous la forme d'une «constante cosmologique». Toutefois, au début des années 1930, Lemaître montra que l'équilibre ainsi obtenu est instable.

repose sur l'étude des concentrations en uranium et en plomb dans ses roches les plus anciennes ; la valeur admise aujourd'hui est d'environ 4,5 milliards d'années.

Le modèle d'un Univers dont le rayon croît de façon exponentielle lève la difficulté en donnant à l'Univers un passé infini durant lequel sa taille est quasi constante, tout en admettant une période d'expansion récente en accord avec l'observation de la fuite des galaxies. Aujourd'hui on attribue à la constante de Hubble une valeur de l'ordre de 70 kilomètres par seconde et par mégaparsec (les estimations des distances des galaxies ont été considérablement améliorées). On en déduit que l'Univers est âgé de quelque 14 milliards d'années, ce qui est compatible avec les données géologiques.

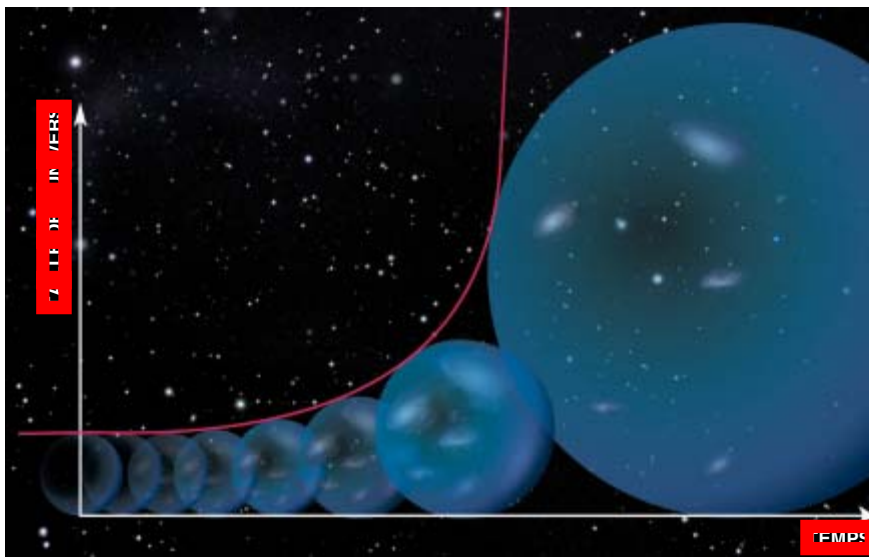
La formation des galaxies

Dans le cadre de ses modèles d'Univers en expansion, Lemaître s'attache alors à décrire un premier scénario de formation des galaxies. Il exploite pour cela les méthodes qu'il a développées durant l'année 1924-1925, lorsqu'il travaillait à l'Institut de technologie du Massachusetts. Cette année-là, Eddington a attiré son attention sur une question soulevée par les travaux de l'astronome allemand Karl Schwarzschild. Ce dernier avait trouvé une solution aux équations d'Einstein décrivant le champ de gravitation à l'intérieur et à l'extérieur d'une boule de matière en supposant que la densité y est homogène (ces hypothèses constituaient un modèle très simplifié d'étoile). La solution obtenue pour l'intérieur de la sphère fait apparaître un paradoxe. En ajoutant de la matière à cette sphère homogène, on fait croître son rayon en même temps que sa masse. Cependant, les calculs montrent qu'au-delà d'une taille (et donc d'une masse) limitée la pression au centre de l'étoile devient infinie. Il semble donc qu'aucun astre ne puisse exister au-delà de cette limite. Toutefois, comme le fit remarquer Eddington, l'hypothèse d'une densité uniforme de matière est peu conforme à l'esprit de la relativité, car, dans cette théorie, la densité de matière n'est pas une grandeur invariante (la masse peut se transformer en énergie et des observateurs en mouvement les uns par rapport aux autres ne constatent pas, pour les mêmes objets, des énergies identiques).

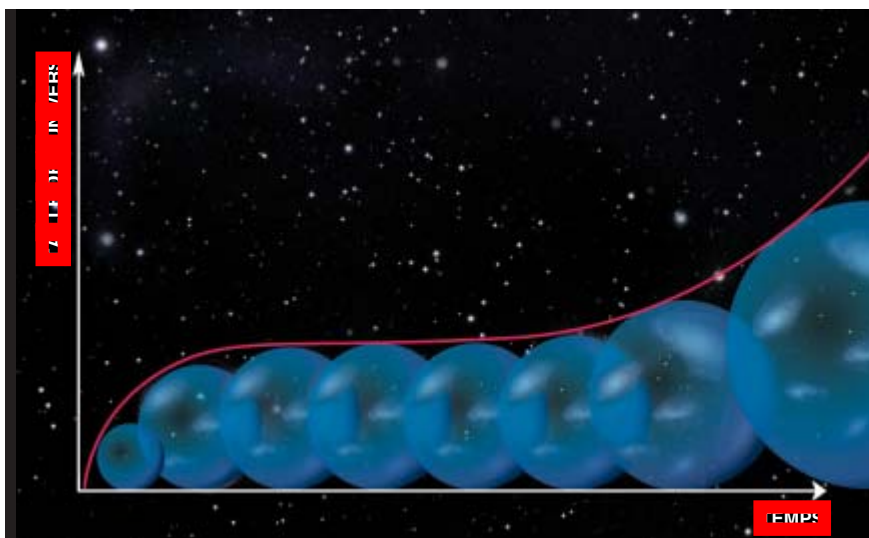
Lemaître refait les calculs de Schwarzschild en abandonnant l'idée d'une densité constante (il ne considère comme constante qu'une grandeur invariante propre à la relativité, la « trace du tenseur énergie-impulsion » qui combine la densité d'énergie et la pression de la matière). Il montre alors, contre l'attente d'Eddington, que le paradoxe découvert par Schwarzschild subsiste : il existe bien un rayon limite

au-delà duquel aucun astre ne peut plus être en équilibre.

L'aspect le plus important de ce travail est qu'il permet à Lemaître d'étudier des espaces à symétrie sphérique remplis d'un fluide dont la densité n'est pas nécessairement homogène. Dans les années 1930, il en tire profit pour proposer un modèle de formation des galaxies à partir de fluctuations locales de la densité de matière dans



4. UN UNIVERS EN EVOLUTION, MAIS ETERNEL, est proposé par Lemaître en 1927. L'expansion de l'espace explique le fait que les galaxies se fuient les unes les autres à une vitesse proportionnelle à la distance qui les sépare. Dans ce modèle, la constante cosmologique n'est pas nulle et la force répulsive qu'elle engendre accélère l'expansion de façon exponentielle. Le commencement de l'Univers est repoussé à l'infini vers le passé où il finit par se « raccorder » au modèle statique d'Einstein. La démonstration qu'un Univers statique est instable conduira Lemaître à abandonner ce dernier modèle.



5. DEUX FORCES ANTAGONISTES se disputent le contrôle de l'Univers dans le modèle que Lemaître défendra jusqu'à la fin de sa vie. Le monde commence par une singularité initiale à une date finie dans le passé, puis son expansion est ralentie par la gravitation. Initialement très faible, la force répulsive engendrée par la constante cosmologique rattrape la gravitation, et l'Univers connaît une phase de quasi-équilibre qui dure plus ou moins longtemps suivant la valeur de ladite constante. Par la suite, la formation des amas de galaxiesrompt l'équilibre et l'expansion reprend, de façon accélérée cette fois. Cette troisième phase correspond à l'époque actuelle.

un Univers en expansion. Selon lui, des particules de matière contenues dans l'Univers se sont agglutinées au hasard à la faveur de fluctuations statistiques. Des zones de densité légèrement supérieure à la moyenne sont apparues. En s'effondrant sous leur propre poids et en attirant la matière environnante, ces fluctuations de densité auraient donné naissance aux galaxies, regroupées par la suite en amas de galaxies, les plus

grandes structures observées jusqu'ici dans l'Univers. Lemaître montre alors que l'estimation de la taille des amas formés selon son modèle est en accord avec les mesures faites par Hubble pour l'amas de galaxies de Coma.

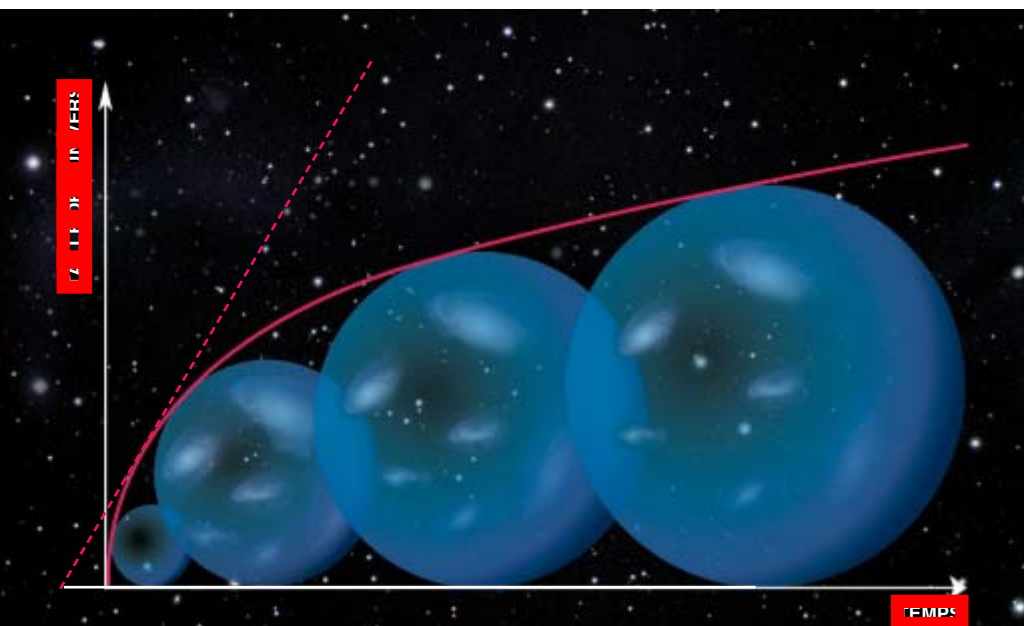
Toutefois, aujourd'hui, on sait que les fluctuations envisagées par Lemaître (des fluctuations statistiques dans un Univers essentiellement homogène) ne suffisent pas à produire les

grandes structures que l'on observe dans l'Univers. Le rayonnement fossile capté par le satellite COBE fournit une image de l'Univers lorsqu'il n'était âgé que de 300 000 ans environ et révèle l'existence de variations de densité dont on pense qu'elles sont à l'origine des grandes structures. Ainsi, l'idée selon laquelle les galaxies et les amas de galaxies se sont formés à partir de condensations locales de matière demeure aujourd'hui celle que privilégient les astronomes. Cependant, beaucoup d'astrophysiciens imaginent qu'il s'agit plutôt de fluctuations microscopiques de nature quantique qui ont été amplifiées au cours de l'inflation de l'Univers (une phase d'expansion exponentielle qui aurait multiplié toutes les distances par un facteur 10^{50} en quelque 10^{-32} seconde).

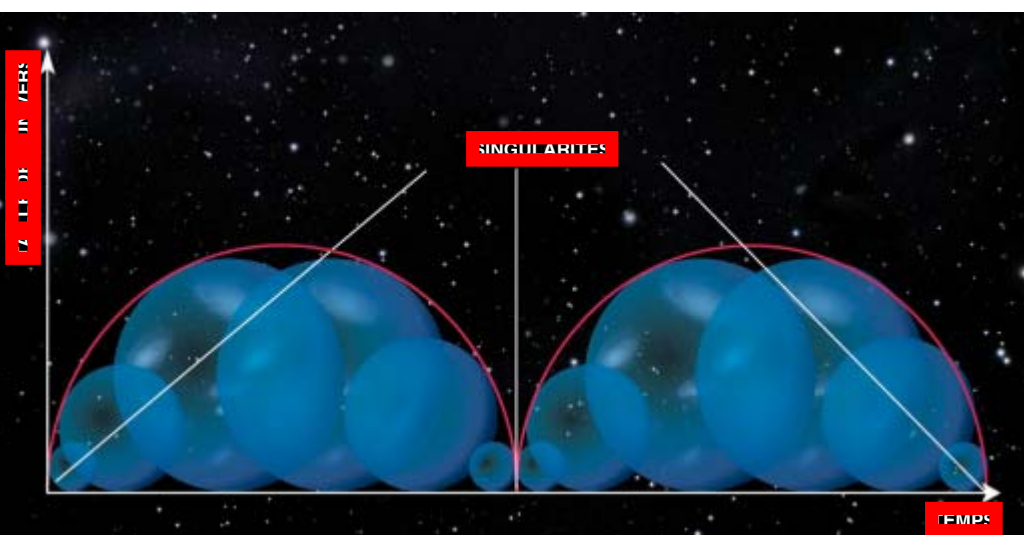
L'âge de l'Univers

Au début des années 1930, Eddington contribue de nouveau à orienter le cours des travaux de son ancien élève en suscitant chez lui un vif intérêt pour la question de l'origine de l'Univers. En effet, en 1931, l'astronome de Cambridge publie un article dans la revue *Nature* où il déclare : «L'idée que le présent ordre des choses ait eu un commencement me répugne philosophiquement.» Lemaître réagit en publiant, dans la même revue, une courte note montrant que la thermodynamique et la théorie quantique peuvent donner un sens physique à un commencement du monde. Dans cet article, il décrit un état initial de l'Univers où tous les quanta d'énergie sont rassemblés en un seul, qu'il nomme «atome primitif». En dehors de lui, les notions d'espace et de temps n'ont pas de sens. Cet état très «ordonné» est instable, et des désintégrations successives comparables aux désintégrations radioactives engendrent progressivement, à partir de l'atome primitif, la matière, l'espace et le temps, tels que nous les connaissons aujourd'hui. Pour Lemaître, «un tel commencement du monde est suffisamment éloigné du présent ordre des choses pour n'être pas répugnant du tout». L'hypothèse annonciatrice de la théorie moderne du Big Bang fait son chemin.

Par ailleurs, au début des années 1930, Lemaître montre que le moindre changement dans la répartition de la densité de matière de l'Univers peut donner à la force répulsive liée à la constante cosmologique l'avantage sur



6. LORSQUE SEULE LA GRAVITATION DETERMINE L'EVOLUTION DE L'UNIVERS, la constante de Hubble suffit à estimer son âge. L'attraction entre les galaxies freine leur fuite réciproque et ralentit l'expansion. Par conséquent, la constante de Hubble, qui n'est autre que la pente de la courbe donnant la taille de l'Univers en fonction du temps, diminue à mesure que l'Univers vieillit. On calcule facilement l'âge qu'aurait l'Univers si son expansion avait été constante et caractérisée par la valeur actuelle de la constante de Hubble (ligne pointillée). Cet âge théorique est une borne supérieure pour l'âge réel de l'Univers.



7. L'UNIVERS PHOENIX enflé, atteint sa taille maximale et s'effondre sur lui-même avant de recommencer un nouveau cycle d'expansion et de contraction. Il fait apparaître des points de densité infinie, des singularités, qui semblent ne pas avoir de sens physique. Einstein se demandait si, en admettant que l'Univers n'est pas rigoureusement isotrope, les calculs ne feraient pas disparaître ces singularités. Lemaître montre que ce n'est pas le cas.

la gravitation. L'équilibre entre ces deux forces sur lequel repose le modèle statique d'Einstein est, par conséquent, instable et la formation des galaxies à partir des fluctuations de matière de l'Univers doit suffire à le rompre. Ainsi, il semble impossible que l'Univers d'Einstein reste indéfiniment statique.

Chez Lemaître, l'hypothèse de l'atome primitif est avant tout une intuition physique plus qu'une théorie rigoureusement étayée. Il correspond au choix d'un nouveau modèle cosmologique où l'âge de l'Univers est fini. Il s'agit d'un Univers homogène, une hypersphère dont l'évolution commence par un état de densité infinie, une singularité, que Lemaître considère comme la limite donnée par les lois classiques de la relativité lorsqu'on s'approche des conditions exotiques qui règnent à l'époque de l'atome primitif. L'évolution de l'Univers est dominée par deux forces, la gravitation et la «force répulsive», dont l'intensité est fixée par la constante cosmologique. Ce modèle, qu'il défendra jusqu'à la fin de sa vie, comprend trois phases caractéristiques (voir la figure 5).

De l'importance d'une constante

Au cours de la première phase, qui commence à la singularité initiale, l'Univers entre en expansion et l'espace se remplit des produits de désintégration de l'atome primitif. L'attraction gravitationnelle qui s'exerce entre les particules de matière freine progressivement l'expansion. La deuxième phase correspond à un équilibre entre la gravitation et la force répulsive liée à la constante cosmologique : le rayon de l'Univers reste, momentanément, quasi constant, comme dans l'Univers d'Einstein. La troisième et dernière phase de l'histoire de l'Univers selon Lemaître inclut l'époque actuelle et commence lorsque la formation des grandes structures et des galaxies rompt l'équilibre de la période quasi statique et provoque la reprise de l'expansion accélérée sous l'effet de la constante cosmologique.

Lemaître attache beaucoup d'importance à ladite constante. Il s'oppose en cela à Einstein qui y a renoncé («c'est la plus grande erreur de ma vie», aurait-il déclaré) en même temps que la découverte de la fuite des galaxies l'a obligé à abandonner son modèle d'Univers statique. Visionnaire, Lemaître



3. ROBERT MILLIKAN, GEORGES LEMAITRE ET ALBERT EINSTEIN, en janvier 1933. Après sa rencontre avec Millikan, Lemaître se persuade que les rayons cosmiques sont des reliquats de la désintégration de l'atome primitif.

estime que la mécanique quantique pourrait, un jour, donner un sens physique à cette constante qui semble signifier que le vide est doté d'une certaine densité d'énergie. On pense aujourd'hui qu'Einstein a traité la constante cosmologique avec trop de légèreté et qu'elle n'est pas une simple «option» mais correspond à une caractéristique fondamentale de sa théorie due à l'existence d'une énergie du vide quantique : la mécanique quantique impose que même dans un espace vide des paires de particules et d'antiparticules apparaissent et disparaissent sans cesse (sinon, la valeur de tous les paramètres physiques serait connue avec une précision parfaite – ils seraient tous nuls –, ce qui contredit le principe d'indétermination de Heisenberg).

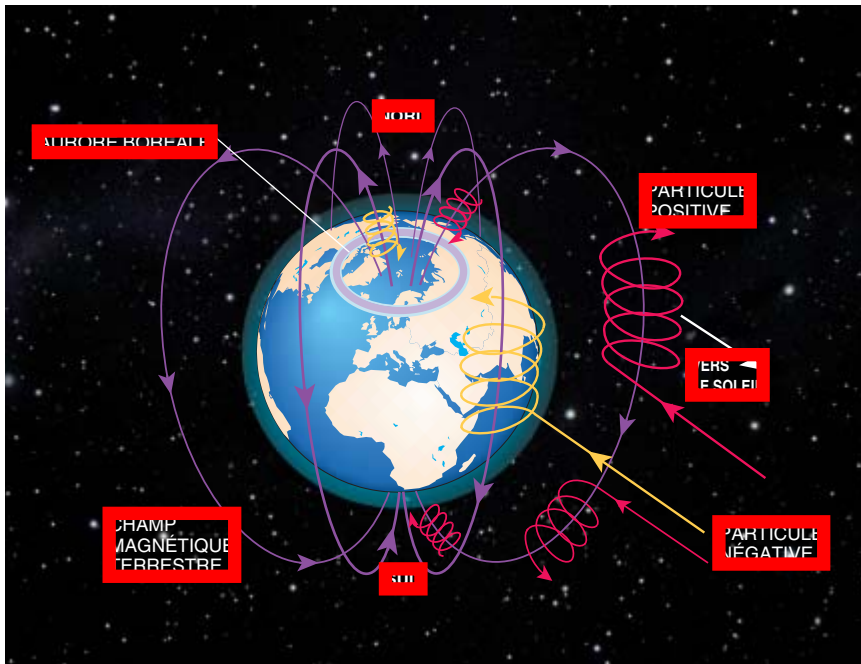
On a effectivement établi récemment, que la constante cosmologique n'est vraisemblablement pas nulle. Par ailleurs, Lemaître insiste sur le fait qu'en changeant la valeur de la constante cosmologique, on modifie l'âge de l'Univers. Selon lui, cela constitue un argument supplémentaire pour la conserver. En effet, si l'évolution de l'Univers n'est gouvernée que par la gravitation, la seule constante de Hubble suffit à déterminer son âge (voir la figure 6). Or, avec la valeur de la constante de Hubble dont on dispose dans les années 1930, cet âge est toujours inférieur à celui du Système solaire et il convient de corriger cette anomalie. Dès 1931, en se fondant sur une borne supérieure de la constante cosmologique, Lemaître évalue à dix milliards d'années l'âge de l'Univers (un bon ordre de grandeur aujourd'hui encore).

L'hypothèse de l'atome primitif et son scénario de synthèse des éléments chimiques par désintégrations successives est complètement abandonné de nos jours. De plus, la géométrie spatiale de l'Univers semble être euclidienne et non pas sphérique, comme Lemaître l'envisagea toujours. Toutefois, diverses observations faites récemment sur les supernovae lointaines semblent plaider pour une valeur non nulle (et positive) de la constante cosmologique et, par conséquent, pour une évolution de l'Univers caractérisée par les trois phases proposées par Lemaître.

Le problème des singularités

Au début des années 1930, on admettait que des galaxies existent en dehors de la nôtre et que l'Univers est en expansion ; on en déduit des modèles de cosmos en évolution permanente, peut-être depuis une durée finie. Cette idée continue cependant à heurter les choix philosophiques de nombreux physiciens. Non seulement il leur faut accepter que l'Univers a eu un commencement, mais de plus, ce commencement se traduit par un état de densité infinie, par une singularité où les équations de la physique perdent tout leur sens.

Ainsi, en janvier 1933, Einstein, qui vient de quitter l'Allemagne pour les États-Unis, rencontre Lemaître à l'Institut de technologie de Californie. Il lui demande si l'on peut éliminer les singularités qui apparaissent dans l'histoire de «l'Univers phoenix», un modèle stipulant que l'Univers enfle à partir d'une singularité, atteint une taille maximale, s'effondre en une nou-



9. LES AURORES BOREALES sont produites par les particules électriquement chargées émises par le Soleil et piégées par le champ magnétique terrestre. Lemaître étudie les rayons cosmiques en supposant qu'ils adoptent des trajectoires semblables lorsqu'ils arrivent au voisinage de notre planète. Les particules chargées positivement (en rouge) circulent autour des lignes de forces du champ magnétique terrestre (violet) en sens inverse des particules chargées négativement (jaune). De ce fait, elles arrivent à l'observateur au sol selon des directions proches de l'Ouest géomagnétique (alors que les particules négatives arrivent de l'Est). Comme cette direction est aussi celle d'où l'on reçoit le plus de rayons cosmiques, Lemaître en déduit qu'ils sont constitués de particules de charge positive.

elle singularité avant de recommencer un nouveau cycle d'expansion et de contraction (voir la figure 7). Comme tous les modèles envisagés par les astronomes, celui-ci est homogène et isotrope (ses propriétés sont les mêmes en tout point de l'espace et dans toutes les directions). Selon Einstein, cette isotropie est peut-être la cause de l'apparition des singularités. Si l'on admettait une légère anisotropie de l'Univers, un Univers en expansion dans deux directions de l'espace et en contraction selon la troisième, par exemple, peut-être l'Univers ne s'effondrerait-il pas complètement sur lui-même et la singularité serait-elle évitée. Lemaître prouve rapidement, à l'aide d'un cas particulier, que la singularité ne disparaît pas, même lorsque l'Univers n'est pas isotrope. Le passage par un état de « rayon » nul semble bel et bien un passage obligé pour la plupart des modèles d'Univers. Cette « expérience mathématique » préfigure les théorèmes sur les singularités de Roger Penrose et Stephen Hawking qui montreront, par des méthodes globales, que des singularités physiques apparaissent de façon inévitable dans beaucoup de modèles possibles d'espace-temps.

Les rayons cosmiques

Selon Lemaître, une des conséquences de l'hypothèse de l'atome primitif est l'existence de particules chargées, hautement énergétiques, produites lors des premières désintégrations de cet atome. Après une rencontre avec Robert Millikan, Lemaître se persuade que ces particules ne sont autres que les rayons cosmiques que l'on capte à haute altitude et dont on ignore encore, au début des années 1930, la nature précise et l'origine. La détection et l'étude des rayons cosmiques présente pour lui une importance cruciale, car il s'agit, selon ses propres termes, de « hiéroglyphes » qu'il faut déchiffrer pour connaître les tout premiers instants de l'Univers.

Lemaître et son collègue de l'Institut de technologie du Massachusetts, Manuel Sandoval Vallarta, entreprennent d'étudier les propriétés des trajectoires des rayons cosmiques. Plusieurs observateurs ont montré que l'intensité du rayonnement cosmique varie suivant la latitude géomagnétique (la latitude mesurée non pas à partir du pôle géographique de la Terre mais à partir de son pôle magnétique). Forts de leur hypothèse selon laquelle ces rayons sont des particules chargées

issues de la désintégration de l'atome primitif, Lemaître et Vallarta s'attachent à calculer l'interaction entre ces particules et le champ magnétique terrestre pour expliquer, notamment, l'« effet de latitude ».

En 1907, un problème semblable avait déjà été abordé par Carl Störmer de l'Université d'Oslo. Störmer étudiait l'interaction des particules chargées émises par le Soleil et du champ magnétique terrestre afin d'obtenir une théorie complète des aurores boréales. Ces phénomènes lumineux sont produits par des particules chargées émises par le Soleil. En s'approchant de la Terre, elles adoptent une trajectoire en hélice le long des lignes du champ magnétique terrestre et sont canalisées vers les pôles. Lorsqu'elles pénètrent dans l'atmosphère, elles ionisent les atomes d'azote et d'oxygène de l'air qui, après leur passage, se désexcitent en émettant de la lumière.

En supposant que les rayons cosmiques – des particules beaucoup plus énergétiques que celles du vent solaire – proviennent de toutes les directions de l'espace (et non plus seulement du Soleil), Lemaître et Vallarta complètent l'approche de Störmer pour étudier leurs trajectoires dans le champ magnétique terrestre. À l'Institut de technologie du Massachusetts, à l'aide d'un ordinateur analogique – la « machine de Bush » – capable d'intégrer des systèmes d'équations différentielles et de donner des représentations graphiques des solutions, ils parviennent à dessiner et à étudier des milliers de trajectoires de rayons cosmiques. Grâce à ce travail, Lemaître et Vallarta confirment que la quantité de rayons cosmiques reçue sur Terre doit varier selon latitude géomagnétique. En outre, ils établissent que les particules qui constituent les rayons cosmiques sont surtout des particules chargées positivement.

En recherchant ces « hiéroglyphes » vestiges des débuts « explosifs » de l'Univers, Lemaître est l'un des premiers physiciens à proposer l'existence d'un rayonnement fossile susceptible de donner une base expérimentale à la cosmologie. Bien sûr, le rayonnement fossile, effectivement détecté dans les années 1960, est de nature très différente de celui dont Lemaître avait postulé l'existence, et l'on considère aujourd'hui que les rayons cosmiques sont des particules très énergétiques – des protons et des noyaux atomiques

égères – produites bien plus tard dans l'histoire de l'Univers, notamment par les supernovae. Toutefois, certaines particules dont l'énergie est de l'ordre de 10^{19} électron-volts ne peuvent pas être produites de cette façon, et certains physiciens pensent qu'elles proviennent de la désintégration de particules massives et exotiques, produites dans les tout premiers instants de l'Univers. En 1998, un partisan de cette théorie, Michael Hillas de l'Université de Leeds, en Grande-Bretagne, concluait ainsi un de ses articles : « Il est tout à fait possible que Lemaître ait été loin d'avoir tort. »

Lemaître possédait une passion pour le calcul numérique et pour les machines à calculer. À Louvain, il acquit toute une série de machines à calculer mécaniques, puis électromécaniques. En 1958, il fit venir le premier ordinateur de l'Université catholique de Louvain, un Burroughs E101, qu'il utilisa pour effectuer des calculs relatifs aux modèles d'amas de galaxies développés dans les années 1940 et 1950 à la suite de ses travaux sur les condensations de matière dans l'Univers en expansion.

Des mathématiques à la cosmologie

Pionnier de la cosmologie, Lemaître fut également, comme beaucoup des théoriciens de la physique classique, un mathématicien de premier plan qui apporta des solutions originales à certains problèmes de mécanique céleste (problème des trois corps) ou encore de calcul numérique (transformée de Fourier rapide). À son époque, d'autres, plus intéressés par les théories prometteuses de la mécanique quantique, se détournaient de ce genre de problèmes, tandis que les mathé-

maticiens s'intéressaient plutôt aux structures abstraites dans le style boussinesquien. Dans les milieux mathématiques de la fin des années 1950, il faisait figure de marginal. Sa passion pour les ordinateurs, pour le calcul numérique et l'expérimentation mathématique était pourtant très moderne à l'instar de ses idées sur l'histoire de l'Univers – qu'il défendit parfois contre Einstein lui-même. La découverte récente de l'importance de la constante cosmologique confirme son intuition et constitue peut-être le plus grand défi lancé aux physiciens pour le siècle qui s'ouvre.

Dominique LAMBERT, docteur en sciences physiques et en philosophie de l'Université catholique de Louvain, est chargé de cours de philosophie et d'histoire des sciences aux Facultés universitaires Notre-Dame de la Paix à Namur.

D. GODART et M. HELLER, *Cosmology of Lemaître*, in *History of Astronomy Series*, Pachart Publishing House, Tucson, vol. 3, 1985.

H. KRAGH, *Cosmology and Controversy*

The Historical Development of Two Theories of the Universe, Princeton University Press, 1996.

A. FRIEDMAN et G. LEMAÎTRE, *Essai de cosmologie*, précédé de *L'invention du Big Bang* par J.-P. LUMINET (textes choisis, présentés, traduits et annotés par J.-P. Luminet et A. Grib), Éditions du Seuil, Collection Sources du Savoir, 1997.

D. LAMBERT, *Un atome d'Univers. La vie et l'œuvre de Georges Lemaître*, Éditions Racine/Éditions Lessius, 2000.

L'Inserm est le premier organisme de recherche français qui se consacre à la santé. Il rassemble plus de 10 000 personnes dans ses 300 laboratoires. Disposant d'un budget de plus de 450 millions d'euros, l'Inserm collabore avec plus de 93 pays et 260 industriels.

Inserm

Institut national
de la santé et de la recherche médicale

Dans le but de soutenir les jeunes chercheurs et de renforcer la dynamique scientifique de la recherche biomédicale, clinique et en santé publique,

l'Inserm lance **Avenir**, un appel à projets pour post-doctorants et jeunes chercheurs titulaires

Ce programme permettra aux jeunes scientifiques retenus de bénéficier :

- d'une aide financière d'au moins 60 000 euros par an pendant 3 ans
- d'un espace équipé d'environ 50 m²
- de l'accès aux plates-formes technologiques
- de la possibilité de composer son équipe grâce, notamment, à l'accueil d'un post-doctorant étranger
- d'une allocation pour les post-doctorants de 2 300 euros net par mois.

Il s'adresse aux jeunes chercheurs, qu'ils soient titularisés ou post-doctorants français et étrangers, présentant un projet scientifique de haut niveau.

45 projets pourront être sélectionnés.

Pour obtenir les informations complètes sur l'appel d'offres **Avenir 2001**, connectez-vous sur : Inserm.fr ou appelez au : 01 44 23 67 05

ou adressez un mail à : postel-vinay@tolbiac.inserm.fr

Date limite de dépôt des dossiers : 29 octobre 2001

Histoire démographique et génétique du Québec

ALAIN GAGNON • HELENE VEZINA • BERNARD BRAIS

Grâce aux registres paroissiaux anciens, les démographes ont reconstitué la généalogie des Québécois d'ascendance française. D'après ces données historiques, les généticiens expliquent pourquoi certaines maladies héréditaires sont fréquentes dans l'Est québécois, alors qu'elles sont rarissimes dans le reste du monde.

« Il serait utile de partir de l'étude des faits et des hommes et de revivifier l'histoire dans ces observations curieuses et fécondes que présente l'étude des populations, des familles elles-mêmes, qui sont la trame essentielle des sociétés que l'on décrit. [...] Il faut utiliser largement cette source historique constituée par les registres paroissiaux et savoir en tirer toutes les données qui en proviennent ». Dès 1860, au cours d'un voyage au Canada, l'historien français Rameau de Saint-Père avait souligné l'intérêt des registres paroissiaux en démographie historique. La tenue de ces registres de baptêmes, de mariages et de sépultures remonte au début de la colonisation. En suivant à la trace la progression du peuplement sur le territoire, les prêtres ont scrupuleusement consigné par écrit la naissance, le mariage et le décès des premiers Canadiens et de leurs descendants.

Grâce à ces registres paroissiaux, les Québécois savent d'où venaient les fondateurs de la population, combien ils étaient où ils se sont installés dans le vaste territoire qu'on appelait alors la Nouvelle-France, combien ils ont eu d'enfants... Autant de questions qui demeurent sans réponse pour la quasi-totalité des autres populations humaines. En 1966, voulant exploiter le potentiel des archives québécoises, les démographes de l'Université de Montréal entreprennent de reconstituer l'ensemble des événements démographiques survenus dans la population d'ascendance française.

Aujourd'hui, nous disposons d'un fichier informatisé de population, nommé *Registre de population du Québec ancien*, couvrant l'ensemble du territoire québécois des origines, en 1608, jusqu'à 1800. Le registre contient environ 712 000 actes de baptême, de mariage et de sépulture, couplés de telle sorte que la généalogie de tous les individus mariés dans la colonie avant 1800 peut être reconstituée, et ce jusqu'aux fondateurs de la population.



En 1971, la construction d'un autre fichier de population est entreprise. Ce fichier, nommé BALZAC, contient aujourd'hui plus de 670 000 actes et couvre les régions du Saguenay-Lac-Saint-Jean et de Charlevoix, situées au Nord-Est de la ville de Québec (voir la figure 4). En voie d'extension à l'ensemble du Québec pour les XIX^e et XX^e siècles, le fichier comprend en outre de nombreuses généalogies constituées de plus de 265 000 fiches individuelles.

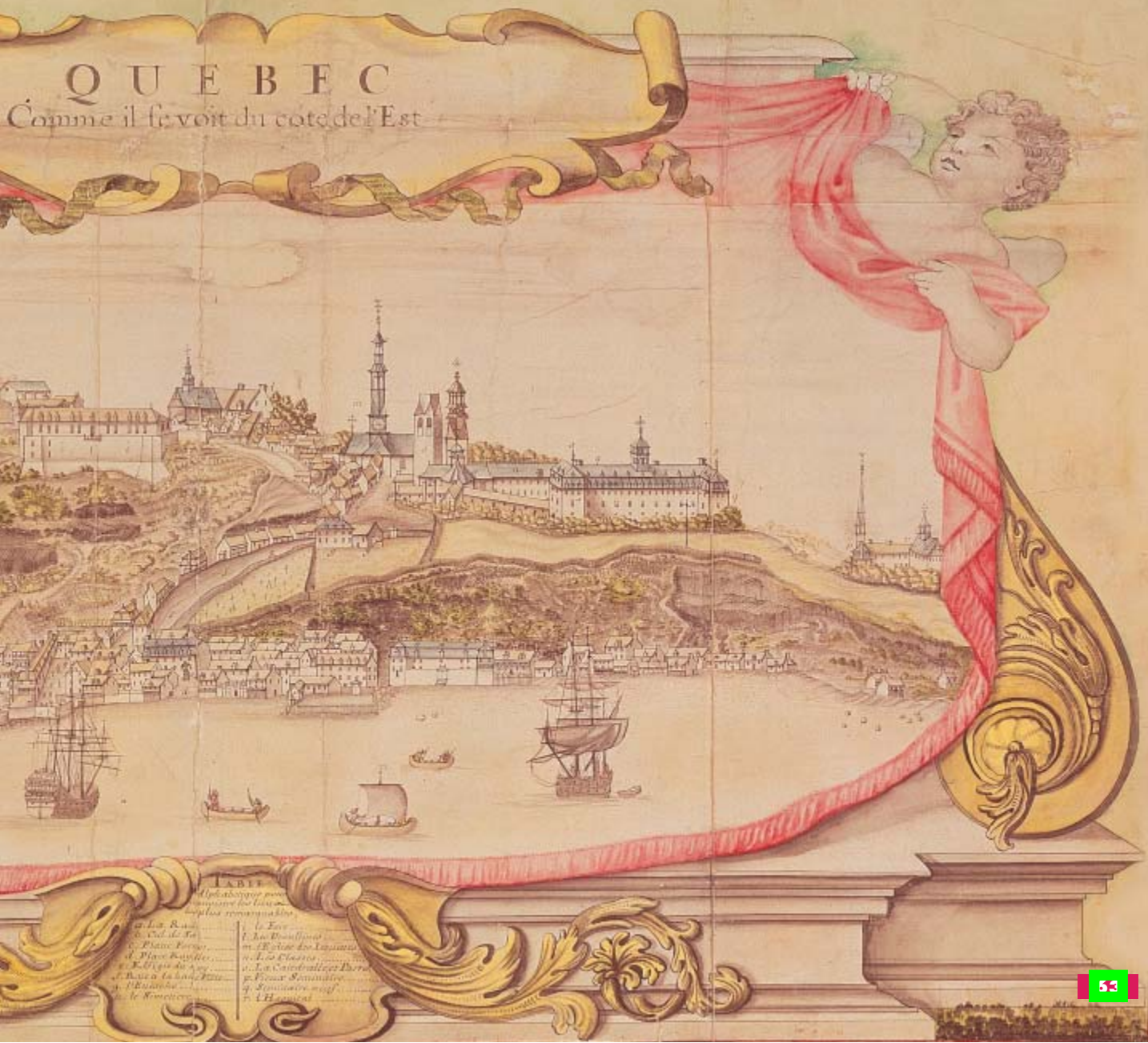
Créés à des fins d'études démographiques, ces deux outils ont aussi été utilisés par des spécialistes des sciences humaines, tels des historiens, des linguistes et des anthropologues, et par des spécialistes des sciences de la vie, tels des biologistes, des médecins et des généticiens. En effet, dans l'Est du Québec, on trouve un certain nombre de maladies héréditaires rares, voire inexistantes ailleurs dans le monde. D'où viennent ces maladies? Comment ont-elles atteint une fréquence si élevée? Pourquoi sont-elles différentes de celles que l'on trouve dans le reste du pays? Les généticiens ont mené des enquêtes sur la généalogie d'individus vivants atteints de ces maladies héréditaires afin

d'expliquer les causes de leur répartition sur le territoire. Ainsi dispose-t-on d'outils pour étudier la constitution et l'évolution de la population québécoise et pour en déduire des informations en génétique et en épidémiologie. Le Québec est devenu, au même titre que la Finlande ou l'Islande, un laboratoire d'étude de la génétique des populations. Dans l'histoire démographique du Québec, les conditions initiales ont été déterminantes. Remontons aux origines du peuplement du Québec.

De l'Ancienne à la Nouvelle-France

Entre 1534 et 1541, le navigateur Jacques Cartier effectue trois voyages en Amérique, cherchant un passage vers l'Asie par le Nord du Nouveau Monde, et prend possession

L. LA VILLE DE QUEBEC EN 1688, 80 ans après sa fondation. Il s'agit d'un détail de la carte de l'Amérique septentrionale, commandée par Louis XIV, désireux de connaître «les frontières entre la Nouvelle-France et la Nouvelle-Angleterre».



d'un vaste territoire au nom de François 1^{er}. Il rencontre les Amérindiens, explore le Saint-Laurent et projette d'établir sur ses berges une colonie française, en faisant miroiter à la couronne les grandes richesses qu'il prétend avoir découvertes. Pourtant, l'entreprise qui donnera naissance au Canada français ne voit le jour qu'en 1608 : mandaté par Henri IV, l'explorateur Samuel de Champlain, qui deviendra plus tard lieutenant-gouverneur du Canada, remonte le Saint-Laurent et fonde la ville de Québec dont le nom signifie «là où la rivière se rétrécit» dans la langue algonquienne (celle d'une tribu indienne du Canada). La France mercantiliste cherche à asseoir ses prétentions territoriales en Amérique afin d'y développer le commerce des fourrures mais les premiers hivers sont rudes. Plusieurs pionniers n'y survivent pas sans compter qu'ils subissent les assauts des Anglais et de leurs alliés iroquois. Pour ces raisons, la jeune colonie n'a pas bonne réputation auprès des Français, qui sont peu nombreux à émigrer. De surcroît, l'or et les métaux précieux, promis un siècle plus tôt par Jacques Cartier, brillent toujours par leur absence.

Quand Louis XVI entreprend de consolider l'établissement de la Nouvelle-France en 1663, soit 55 ans après sa fondation, à peine 2 500 colons y

sont installés. Le roi s'en inquiète et ordonne que d'autres y soient envoyés, notamment les Filles du Roi (des orphelines, dont la plupart vivaient à l'Hôpital de la Salpêtrière pour rétablir l'équilibre des sexes ; il n'y a alors dans la colonie qu'une femme à marier pour six ou sept hommes. À la même époque, le régiment Carignan est dépêché sur le territoire pour contrer les menaces anglaises et iroquoises. Au total, entre 1663 et 1680, 1 968 pionniers s'établissent dans la colonie, alors qu'entre 1608 et 1663, on n'en compte que 412. Grâce au *Registre de population du Québec ancien*, réalisé par Hubert Charbonneau, Jacques Légaré et leurs collègues de l'Université de Montréal dans le cadre du *Programme de recherches en démographie historique*, on connaît avec précision le lieu de provenance de la plupart de ces pionniers (voir la figure 3).

Toutefois, cet effort de peuplement n'est que de courte durée : après 1680, le nombre de nouveaux arrivants diminue. D'ailleurs, dès 1666, le ministre Jean-Baptiste Colbert rappelle à l'intendant Jean Talon qu'«il ne serait pas de la prudence [du Roi] de dépeupler son Royaume comme il faudrait faire pour peupler le Canada». À l'instar de ses contemporains, il était convaincu que la population française diminuait. Pour que celle de la Nouvelle-France augmente, on devait donner la primauté à l'accroissement naturel, c'est-à-dire à la reproduction. En 1760, le Canada est conquis par l'Angleterre. À partir de ce moment, les visées de Colbert se réalisent, puisque les immigrants français se raréfient et que la fécondité des

Canadiens français atteint des sommets : jusqu'à la fin du XIX^e siècle, le nombre moyen d'enfants par femme demeure supérieur à sept.

Une répartition contrastée des maladies héréditaires

Au cours des siècles, et surtout depuis quelques décennies, la population s'est enrichie de nouveaux immigrants de provenances diverses. Cependant, ils n'ont apporté qu'une faible contribution au patrimoine génétique canadien français, parce qu'ils étaient peu nombreux par rapport à la population de souche plus ancienne qui a continué à augmenter rapidement. D'après les résultats du *Programme de recherches en démographie historique* de l'Université de Montréal, parmi les 3 380 pionniers arrivés avant 1680, 2 600 sont à l'origine des deux tiers du patrimoine génétique des six millions de Canadiens français actuels (qui forment plus de 80 pour cent de la population québécoise). De plus, seulement 575 de ces pionniers ont fourni le tiers de ce même patrimoine génétique. Ainsi, une grande partie du génome canadien français provient d'un nombre réduit de fondateurs.

Le modèle de «l'effet fondateur» décrit bien la naissance de la population québécoise. On parle d'effet fondateur lorsqu'un groupe restreint d'immigrants s'établit sur un nouveau territoire et donne naissance à une nouvelle population. Par la suite, c'est le phénomène de «dérive génétique» qui prend le relais : certains gènes se diffusent en grand nombre alors que d'autres disparaissent. À mesure que

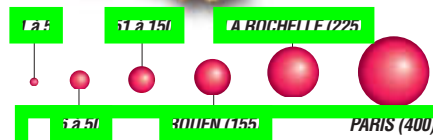
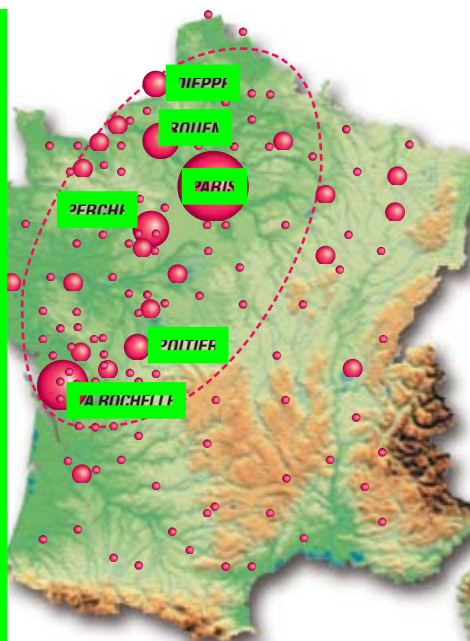
2. LE NAVIGATEUR FRANÇAIS Jacques Cartier, parti de Saint-Malo pour trouver un passage vers l'Asie par le Nord du Nouveau Monde, aborde au Canada en 1534 et prend possession d'un vaste territoire au nom de François 1^{er}. Cette peinture, où Cartier (à droite) et ses compagnons sont représentés, date de la même époque (entre 1536 et 1542).



es générations se succèdent, un gène défectueux, mais rare dans la population mère, peut devenir très fréquent dans la population qui en est issue. Aussi, au bout de plusieurs générations, beaucoup de descendants partagent un lot de gènes identiques, reçu de leurs ancêtres communs. Selon Bertrand Desjardins, démographe à l'Université de Montréal, deux individus d'ascendance canadienne française, choisis au hasard dans la population, partagent au moins un ancêtre arrivé sous le Régime français (avant 1760).

Selon les travaux de l'Institut interuniversitaire de recherches sur les populations, en collaboration notamment avec le Réseau de médecine génétique appliquée du Québec, les conséquences de l'effet fondateur qui se caractérisent par une répartition particulière des maladies héréditaires, se sont surtout manifestées dans l'Est (voir la figure 4). Globalement, les maladies héréditaires dépistées en avant de la ville de Québec (surtout au Saguenay et à Charlevoix) sont en petit nombre, mais ont une fréquence élevée. Certaines de ces maladies n'existent nulle part ailleurs. C'est le cas de l'ataxie spastique, qui provoque une dégradation progressive de la moelle épinière. La fréquence du gène mutant qui en est responsable est d'environ cinq pour cent au Saguenay et à Charlevoix. Le gène mutant responsable de la tyrosinémie, maladie entraînant une cirrhose du foie dès les premiers mois de la vie, a également une fréquence élevée, alors qu'il est très rare à l'extérieur de ces régions. En revanche, la phénylcétonurie (dont les symptômes sont un grave retard mental, des troubles de la pigmentation, de l'eczéma et des convulsions) et la dystrophie musculaire facio-scapulo-humérale (correspondant à un relâchement des muscles du visage de l'épaule et du bras), des maladies héréditaires relativement fréquentes, sont presque absentes de ces régions.

Au centre et à l'Ouest du Québec c'est le contraire : dans une aire délimitée par les villes de Québec, Sherbrooke, Montréal et Hull, on observe une large éventail de maladies héréditaires dont chacune est faiblement représentée (elles touchent une personne sur 100 000 ou moins selon les maladies). Par exemple, on rencontre la mucoviscidose (qui cause notam-



3. LA FRANCE a massivement contribué au patrimoine génétique québécois. Parmi les 3 380 pionniers arrivés avant 1680, 99 pour cent étaient français. Cette carte représente les lieux de provenance des fondateurs. Les régions situées à l'intérieur d'un ovale délimité par les villes de La Rochelle, Poitiers, Paris, Rouen et Dieppe ont fourni le plus grand nombre d'immigrants. Une trentaine d'immigrants provenant d'autres pays d'Europe et une douzaine d'Amérindiens déjà sur place ont aussi participé au peuplement initial.

ment une insuffisance respiratoire grave), l'ataxie de Friedreich (une maladie neurodégénérative se manifestant par une mauvaise coordination des mouvements, des troubles moteurs et une faiblesse musculaire) ou la phénylcétonurie. Cette répartition est caractéristique des populations européennes ou d'origine européenne, telle la population des États-Unis. Elle ne découle pas d'un effet fondateur, contrairement à la répartition observée dans l'Est. Toutes ces données indiquent une fragmentation spatiale du génome québécois.

Pour étudier cette fragmentation, nous avons analysé la contribution génétique des fondateurs pour chaque région du Québec ancien. La contribution génétique désigne la part probable qu'occupe un fondateur dans le génome d'un individu ou d'une population. Par exemple, la contribution d'un grand-père à son petit-fils est égale à 0,25 : en moyenne, le quart des gènes

est au hasard chez le petit-fils proviennent du grand-père, puisqu'il a quatre grands-parents. Pour estimer la contribution d'un fondateur à une région donnée, on fait la somme de toutes ses contributions à l'ensemble de ses descendants et on divise le résultat par l'effectif total de la région. Grâce au *Registre de la population du Québec ancien*, on connaît la descendance par région de tous les fondateurs jusqu'à 1800, de sorte que l'on calcule aisément la contribution des fondateurs.

Connaissant la répartition spatiale de la contribution des fondateurs, nous avons estimé les proximités génétiques entre les régions : plus deux régions ont des fondateurs communs qui ont eu des contributions comparables, plus elles sont proches génétiquement. Nos calculs des proximités génétiques entre régions, pour la période 1780-1799, montrent l'existence de trois grands regroupements de régions : l'Ouest, le Centre et l'Est (voir la figure 5). Les trois regroupements territoriaux correspondent approximativement à ce que les autorités coloniales nommaient les « Gouvernements ». Dans le regroupement de l'Est, on ne retrouve pas la région du Saguenay-Lac-Saint-Jean, car elle n'a été peuplée qu'à partir de la seconde moitié du XIX^e siècle par un mouvement migratoire provenant principalement de la région de Charlevoix (incluse dans l'analyse).

Pourquoi cette division tripartite ? Parce qu'il y avait trois ports d'entrée en Nouvelle-France, autour desquels la colonie s'est développée : Québec, Trois-Rivières et Montréal. Ainsi, la fragmentation spatiale du génome québécois est due aux circonstances initiales du peuplement.

Un profil génétique particulier : l'Est du Québec

La généalogie de 706 individus atteints de l'une ou l'autre des maladies les plus répandues dans la région du Saguenay-Lac-Saint-Jean (l'ataxie spastique et la tyrosinémie, notamment) confirme que, du point de vue génétique, l'Est du Québec se distingue du centre et de l'Ouest. On a montré que 95 pour cent des ancêtres qui composent pour beaucoup (dont la contribution génétique était élevée) dans l'ascendance des individus atteints



Il s'agit d'une maladie héréditaire se sont établis dans l'Est ; 90 pour cent de leurs enfants et 80 pour cent de leurs petits-enfants y sont restés. En revanche, les fondateurs qui comptent peu dans l'ascendance des malades ne se sont pas établis dans l'Est qu'une fois sur deux et le tiers de leurs petits-enfants y sont restés. Bref, les fondateurs qui ont contribué le plus au pool génétique des personnes atteintes d'une maladie génétique du Saguenay-Lac-Saint-Jean sont « spécifiques » de l'Est du Québec et ont une descendance réduite au centre et à l'Ouest. Cette analyse explique qu'on trouve des maladies spécifiques à l'Est. Cependant, certaines maladies héréditaires, telle la dystrophie oculopharyngée, ont connu une diffusion plus large (voir l'encadré page 68).

L'Est se distingue par la présence d'un nombre réduit de maladies spécifiques dont les fréquences sont élevées. Comment peut-on expliquer cela à la lumière des données historiques et démographiques? L'Est, plus éloigné des ports d'entrée en Nouvelle-France, a reçu beaucoup moins d'immigrants fondateurs que l'Ouest de sorte que l'éventail de mutations introduites a été moins large. Puisque le nombre de mutations différentes était réduit, quelques-unes d'entre elles ont atteint une fréquence élevée. Pour s'en convaincre, songeons à un jeu de 26 cartes au lieu de 52. Dans

ces conditions, chaque carte revient à paraître deux fois plus souvent entre les mains des joueurs : de la même manière, une mutation particulière aura plus de chances de se retrouver dans le bagage génétique d'un habitant de l'Est.

Cependant, le nombre de fondateurs n'explique pas tous les traits génétiques de la population québécoise. Récemment, les démographes de l'Institut interuniversitaire de recherches sur les populations et du Laboratoire d'anthropologie biologique du musée de l'Homme ont souligné que les comportements sociaux peuvent provoquer une augmentation notable de la fréquence de certains gènes mutés et des maladies héréditaires qui leur sont associées. Ces comportements ont eu une incidence déterminante sur le patrimoine génique des régions de Charlevoix et du Saguenay-Lac-Saint-Jean.

Les facteurs sociaux et démographiques

Les gènes défectueux ou mutants responsables des maladies héréditaires sont le produit d'un « accident » survenant dans les cellules sexuelles. Certaines mutations, quand elles sont avantageuses, c'est-à-dire quand elles favorisent la reproduction de ceux qui les portent, se répandent dans une population. Quand elles sont désa-

vantageuses, elles sont éliminées par la sélection naturelle : les porteurs ne survivent pas ou n'ont pas de descendants. Toutefois, il arrive que des mutations néfastes se transmettent sur plusieurs siècles. C'est le cas, notamment, des mutations correspondant à une maladie récessive, telle la mucoviscidose. Pour qu'un individu soit atteint d'une maladie récessive, il faut que ses deux parents soient porteurs du gène défectueux et que chacun d'eux le lui transmette. Puisque le père et la mère ont chacun une chance sur deux de transmettre le gène défectueux, un enfant sur quatre, en moyenne, est atteint. Trois fois sur quatre, l'enfant n'a aucun symptôme, mais s'il est porteur asymptomatique (deux fois sur trois), il transmet la mutation à la descendance. Les maladies qui dépendent de gènes mutés s'expriment après la période de reproduction peuvent également se transmettre de génération en génération. C'est le cas de certaines formes de la maladie d'Alzheimer ou du cancer du sein qui sont des maladies multifactorielles, c'est-à-dire qui dépendent de facteurs de risques génétiques et environnementaux. Puisque ces gènes n'empêchent pas leurs porteurs de se reproduire, ils ne sont pas éliminés.

Afin de mieux comprendre le rôle des facteurs génétiques impliqués dans la maladie d'Alzheimer, nous avons

reconstitué les généalogies de 205 personnes souffrant de cette maladie et de 205 personnes ne présentant aucun déficit cognitif (les témoins), vivant dans la région du Saguenay. Nous avons supposé que les malades, parce qu'ils partagent des gènes contribuant au développement de la maladie, seraient plus apparentés entre eux qu'aux individus du groupe témoin. Selon cette hypothèse, les individus atteints auraient reçu des gènes de susceptibilité (qui entraînent un risque génétique) de leurs ancêtres communs. Les analyses ont été faites en regroupant les sujets selon qu'ils présentaient la forme précoce ou la forme tardive de la maladie et selon qu'ils étaient porteurs ou non d'un exemplaire particulier (l'allèle $\epsilon 4$) du gène codant l'apolipoprotéine E, qui constitue un facteur de susceptibilité important. Les résultats indiquent que les patients présentant la forme tardive de la maladie et qui sont porteurs de l'allèle $\epsilon 4$ sont plus apparentés entre eux qu'aux témoins : ils ont une probabilité plus grande que les témoins de partager des gènes dont certains pourraient être associés à la maladie. De plus, ces patients sont plus consanguins (un individu est consanguin lorsque son père et sa mère sont apparentés) que les témoins, ce qui pourrait indiquer la participation d'un gène récessif.

Selon le «père» du fichier BAL, SAC, Gérard Bouchard, de l'Université du Québec, à Chicoutimi : «Le hasard biologique est aveugle ; ce sont nos conduites qui le guident.» En effet, la

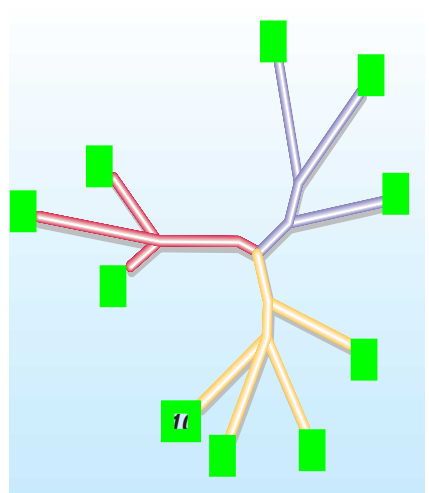
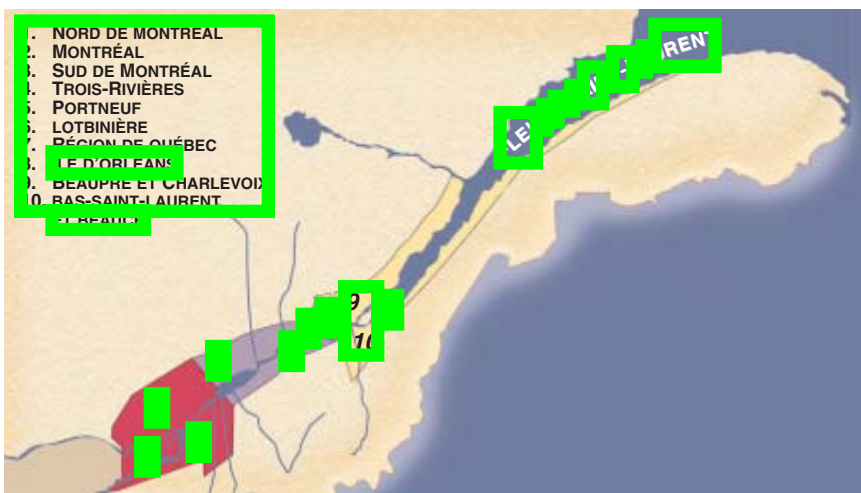
réquence d'un gène mutant, apparu par hasard dans une population, dépend largement des comportements démographiques. Pour rendre compte des comportements, on calcule la «fécondité utile» de chaque individu : c'est-à-dire le nombre de ses enfants qui se sont reproduits dans leur population d'origine. Ainsi, imaginons deux couples : le premier a 12 enfants. Trois demeurent en bas âge, trois émigrent à l'extérieur de la population et deux autres demeurent célibataires. Au total, seuls quatre enfants de cette famille se marient, s'établissent et se reproduisent dans la population. L'autre couple est moins fécond. Il a six enfants, mais tous se marient et s'établissent dans la population avec leurs descendances respectives. Bien qu'il ait été moins fécond, le second couple a plus contribué à la population que le premier. Les variations de fécondité utile d'un couple à l'autre favorisent la diffusion de certains gènes au détriment d'autres et réduisent la diversité génétique.

Les comportements démographiques

Dr, Frédéric Austerlitz et Evelyne Heyer, du Laboratoire d'anthropologie biologique du musée de l'Homme ont montré que ce phénomène ne suffit pas pour expliquer comment certaines mutations (telles celles responsables de l'ataxie spastique et de la tyrosinémie) ont atteint une fréquence de cinq pour cent au Saguenay et à Charlevoix. S'appuyant sur

une étude de Marc Tremblay, de l'Université de Chicoutimi, ils ont mis en évidence le rôle de la corrélation de la fécondité utile : généralement, les individus dont les parents ont eu beaucoup d'enfants utiles ont eu à leur tour beaucoup d'enfants utiles, comme si la faculté de pourvoir à l'établissement d'un grand nombre d'enfants se transmettait de génération en génération. Inversement, les individus issus de petites familles ont eu peu de descendants qui se sont implantés dans leur région. Lorsqu'une grande corrélation de fécondité utile se conjugue à une grande variation de cette fécondité, la réduction de la diversité génétique est importante : non seulement il y a moins de gènes différents qui passent d'une génération à la suivante à cause de la variation de la fécondité utile), mais ce sont surtout les gènes associés aux lignées déjà répandues qui passent (à cause de la corrélation de la fécondité utile). Au Saguenay, la corrélation est de 0,3 (la valeur 0 correspond à une absence de corrélation et la valeur 1 à une corrélation parfaite). Des simulations sur ordinateur montrent qu'une corrélation égale à 0,3 a entraîné une diminution de la diversité génétique d'un facteur 10 en l'espace de 12 générations.

Comment expliquer l'existence de cette corrélation au Saguenay ? Les familles qui venaient de Charlevoix ont bénéficié de l'avantage de l'ancienneté arrivées les premières dans la région du Saguenay, elles ont accaparé plus de terres, où un plus grand nombre



5. L'ANALYSE DE LA CONTRIBUTION GÉNÉTIQUE de 9 659 fondateurs à chacun des 44 000 individus mariés entre 1780 et 1799 révèle l'existence de trois bassins génétiques au Québec (à gauche) : l'Ouest (en rouge ; régions 1, 2 et 3), le centre (en violet ; régions 4, 5 et 6) et l'Est (en jaune ; régions 7, 8, 9 et 10). Ces aires sont celles qui ont été peuplées avant 1800 : à l'époque, la population s'éloignait

peu des rives du Saint-Laurent et de ses principaux affluents (pour plus de lisibilité, on a exagéré la largeur des régions sur la carte). Les régions de 1 à 10 sont d'autant plus proches génétiquement que le chemin à parcourir sur l'arbre des distances génétiques (à droite) pour aller de l'une à l'autre est court. Ainsi, Lotbinière (6) est plus proche génétiquement de la région de Trois-Rivières (4) que de celle de Québec (7).

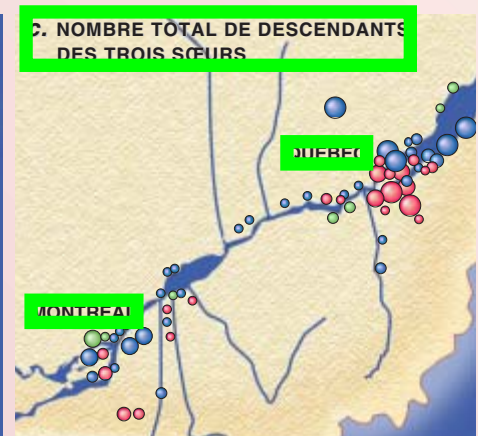
Le suivi d'une mutation particulière

■ Au Québec, grâce aux bases de données généalogiques, on peut reconstituer avec précision l'origine et le destin d'une mutation particulière. Des les années 1960, les médecins généticiens avaient remarqué que la dystrophie oculopharyngée est très fréquente dans les comtés de Montmagny et de l'Islet, situés en aval de Québec, sur la rive Sud du Saint-Laurent.

Cette maladie dominante (une seule copie du gène muté, provenant de l'un ou l'autre des parents, suffit au déclenchement de la maladie) se caractérise, à l'âge adulte, par un relâchement des muscles des paupières, une déglutition difficile, ainsi que par des anomalies de l'œsophage et du pharynx. Les généalogies ascendantes de 89 patients atteints de la maladie dans ces régions convergeaient invariablement vers un couple fondateur, dont l'établissement en Nouvelle-France a eu lieu en 1649. Cela laissait supposer que l'un ou l'autre des membres de ce couple avait introduit la mutation responsable dans la population. Or, en remontant les généalogies d'autres individus atteints, mais vivant dans d'autres régions du Québec ou même aux États-Unis, les démographes de l'Institut interuniversitaire de recherches sur les populations et du Programme de recherches en démographie historique de l'Université de Montréal se sont aperçus que ce couple n'y était pour rien ; étant donné sa nombreuse descendance, il figure dans la généalogie de la plupart des Québécois, dans celle des individus atteints de la dystro-

phie oculopharyngée comme dans celle des autres. En réalité, le couple qui a introduit la mutation n'est jamais venu en Nouvelle-France, mais trois de leurs filles y ont immigré et y ont fait souche en 1648. Puisque nous connaissons l'origine de la plupart des fondateurs, nous connaissons aussi celui de ces trois sœurs : elles sont nées à Niort, dans le Poitou. On peut même retracer la diffusion de la mutation qu'elles ont introduite en Amérique. Quelques cas ont été dépistés aux États-Unis, chez des personnes d'ascendance franco-canadienne. En effet, au XIX^e siècle, environ 900 000 Québécois ont immigré vers les États-Unis pour y chercher du travail : certains des descendants des trois sœurs y ont amené avec eux la mutation, qu'ils ont ensuite transmise à leurs enfants. Au Québec, les cas dépistés se concentrent autour des régions de l'Islet et Montmagny, où la sœur cadette a eu beaucoup de descendants (*représentés par des cercles de couleur bleue*). La benjamine est décédée plus jeune que les deux autres et n'a eu que trois enfants ; c'est pourquoi elle a eu moins de descendants que les deux autres.

La dystrophie myotonique oculopharyngée a été introduite en Amérique par trois sœurs nées à Niort, qui ont émigré au Québec en 1648. Puis, elle s'est diffusée en Amérique de Nord (a). Au Québec, les lieux des personnes atteintes de la dystrophie oculopharyngée (b) correspondent aux endroits où se sont établis les descendants actuels des trois sœurs (c, les plus petits cercles correspondent à 1 descendant, les plus grands représentent 50 à 75 descendants).



de leurs enfants se sont établis. C'est le phénomène des pionniers accapareurs, qui met en application la maxime «premiers arrivés, premiers servis». En marge de ces familles figure une population flottante, celle des migrants, dont les enfants sont moins nombreux à s'établir dans la région que ceux de la vieille souche. Comme l'a montré G. Bouchard, ce phénomène est lié au mode d'occupation du sol sur le front pionnier, où la terre doit être défrichée par la mise en place de vastes réseaux de solidarité, traditionnellement tissés par les liens de parenté et de voisinage, et dont les migrants sont exclus.

Ainsi, la transmission des comportements reproductifs est un phénomène de nature sociale. Des individus très fertiles peuvent transmettre cette disposition à leurs enfants. Cependant, si certains gènes mutés se sont autant répandus au Saguenay, ce n'est pas parce qu'ils étaient avantageux, mais parce qu'ils se sont transmis de génération en génération au sein des familles les plus anciennes de la région qui n'ont cessé de s'agrandir.

Et les autres populations?

Nous supposons que le modèle de la transmission des comportements reproductifs s'applique à d'autres populations fondatrices comme celles des colonies anglaises, espagnoles ou portugaises d'Amérique. Ces populations ont longtemps disposé de grands espaces en friche, ce qui a probablement favorisé l'enracinement à long terme de certaines familles. Le phénomène semble même encore plus général : il a été observé dans une vieille population d'Europe, celle de la Valserine. Dans cette vallée du Haut-Jura, la fréquence de porteurs de la maladie de Rendu-Osler (une maladie rare, qui se manifeste dès l'enfance par des saignements de nez à répétition et par des lésions

de la peau et des muqueuses) est très élevée. L'analyse démographique menée par E. Heyer révèle que la population se compose d'un noyau d'autochtones, entouré des migrants ou des étrangers installés depuis peu dans la population. Ces derniers y ont une descendance réduite, car la majorité de leurs enfants émigrent. Là encore, les causes sont sociales : le nombre de terres étant limité, seuls les autochtones disposent d'un espace cultivable qu'ils transmettent à leurs enfants.

Si la transmission sociale des comportements reproductifs est une caractéristique générale des groupes humains, il faudra réviser bien des modèles de génétique des populations car cette transmission provoque une réduction importante de la diversité génétique. Ainsi, les études démographiques québécoises mettent en évidence ce phénomène, que l'on ne peut observer dans la majorité des populations humaines. Cependant, l'utilisation des données démographiques du Québec n'a pas qu'un intérêt théorique : elle a aussi un intérêt médical.

Grâce aux développements récents de la génétique, on peut aujourd'hui prévoir l'apparition de certaines maladies avant même qu'elles ne se manifestent. C'est là l'objet de la médecine génétique, fondée sur le pronostic plutôt que sur le diagnostic ou sur le traitement des maladies. Dans ce domaine, la connaissance de l'histoire démographique constitue un atout majeur. Comme nous l'avons montré, elle rend compte de la répartition de diverses maladies héréditaires dans certaines populations, et peut ainsi contribuer à évaluer les risques à l'échelle régionale. Au Québec, on peut espérer qu'elle facilitera un jour la mise en place de stratégies de prévention efficaces, dans le respect de la dignité humaine et des droits fondamentaux du citoyen.

Alain GAGNON, démographe de l'Université de Montréal, est chercheur au Laboratoire d'anthropologie biologique du musée de l'Homme. Hélène VEZINA est professeur et chercheur en démographie génétique à l'Université du Québec, à Chicoutimi. Bernard BRAIS est neurogénétiicien au Centre hospitalier de l'Université de Montréal.

Hubert CHARBONNEAU et al., *Naissance d'une population, les Français établis au Canada au XVII^e siècle*, Institut national d'études

démographiques, Presses Universitaires de France et Presses de l'Université de Montréal. Paris et Montréal. 1987.

Gérard BOUCHARD et Marc de BRAEKELEER, *Histoire d'un génome. Population et génétique dans l'Est du Québec*, Presses de l'Université du Québec. Québec. 1991.

Hélène VEZINA et al., *Genealogical Study of Alzheimer Disease in the Saguenay Region of Quebec*, in *Genetic Epidemiology*, n° 16, pp. 412-425. 1999.

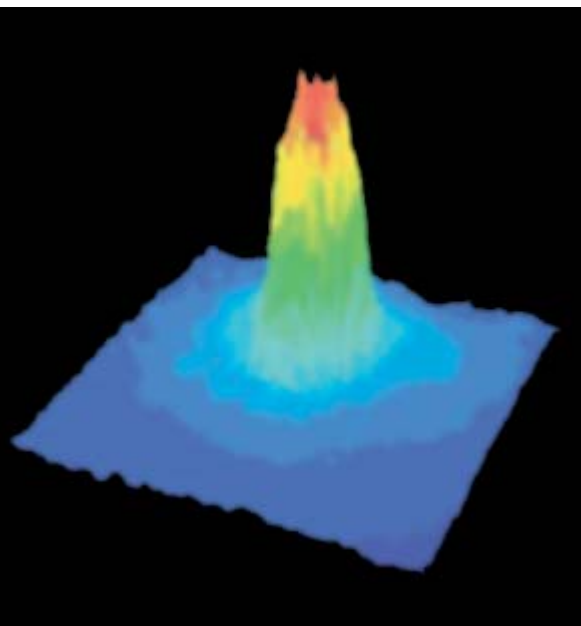
La lumière ralentie

LENE VERSTERGAARD HAU

Des expériences où la lumière est ralentie, voire même arrêtée, ouvrent la voie, notamment, à de nouvelles méthodes de communication optique et à la simulation de trous noirs en laboratoire.

En juillet 1998, dans un avion qui me transportait de Cambridge à Copenhague, je m'émerveillais de voyager plus vite que la lumière.

Depuis cinq mois, à l'Institut Rowland de Cambridge, nous parvenons à ralentir des impulsions lumineuses au-dessous de la vitesse d'un long courrier. Un mois plus tard, nous avions atteint 60 kilomètres par heure... Vers la fin de l'année dernière, nous réussissions à bloquer totalement des impulsions lumineuses à l'intérieur de minuscules nuages de gaz, refroidis au voisinage du zéro absolu, avant de les «relâcher».



1. UN CONDENSAT DE BOSE-EINSTEIN se forme à très basse température, lorsque certains atomes se condensent en un état quantique unique et évoluent de conserve (le pic central indique qu'ici un très grand nombre d'atomes se sont rassemblés pour former un condensat). Cette image représente l'absorption de la lumière par un nuage de sodium refroidi à une température inférieure à 500 milliardièmes de degrés au-dessus du zéro absolu.

Tout lycéen sait que la vitesse de la lumière est l'une des plus inébranlables caractéristiques de l'Univers. Pour autant, dans nos expériences, elle diminue! À l'Institut Rowland, les premières expériences de ralentissement de la lumière réalisées par notre équipe duraient en moyenne 27 heures. Pour ce faire, on fait voyager la lumière au sein d'un milieu constitué d'atomes très froids. Lorsque la température devient suffisamment basse, ces atomes forment un condensat de Bose-Einstein, un système remarquable au sein duquel les atomes sont regroupés en un état quantique unique où ils évoluent de conserve. Un rayon laser dont la fréquence est ajustée contribue à placer l'ensemble des atomes du condensat dans un état qui les rend incapables d'absorber la lumière d'une fréquence donnée : pour une impulsion émise à cette fréquence, le nuage devient transparent, et l'impulsion y progresse très lentement.

Outre son intérêt en physique fondamentale, le ralentissement et l'arrêt de la lumière sont susceptibles d'avoir plusieurs applications. Ainsi, on pourra utiliser de nouvelles méthodes pour sonder les condensats de Bose-Einstein. Par ailleurs, de nouvelles perspectives seront ouvertes dans les domaines des communications optiques, du stockage des données et du traitement quantique de l'information, avec, en ligne de mire, les ordinateurs quantiques qui surpasseront les performances des ordinateurs classiques.

De nombreux matériaux ordinaires ralentissent la lumière. Un rayon lumineux voyage dans l'eau, par exemple, à une vitesse inférieure d'un quart à celle de la lumière dans le vide. Toutefois, le ralentissement de la lumière dans un matériau transparent ordinaire est limité.

La transparence induite

On nomme indice de réfraction le rapport de la vitesse de la lumière dans le vide et de sa vitesse dans un milieu transparent donné. Le diamant, matériau dont l'indice de réfraction est l'un des plus élevés, ne ralentit la lumière que d'un facteur 2,4. Pour obtenir des facteurs de ralentissement de l'ordre de dix millions, on doit recourir à un effet d'origine quantique que l'on nomme la transparence électromagnétique induite, un phénomène observé pour la première fois au début des années 1990, par l'équipe de Stephen Harris, de l'Université de Stanford. La transparence induite permet de rendre un nuage de gaz aussi transparent que du verre pour les rayons lumineux d'une fréquence bien précise. Pour ce faire, on le manipule à l'aide d'un faisceau laser dit de couplage dont la fréquence est elle-même très précisément choisie.

Malheureusement, cette manipulation est perturbée par l'agitation thermique des atomes. En effet, lorsqu'un atome se rapproche d'une source de lumière, il «voit» arriver des photons dont la fréquence est augmentée d'un facteur proportionnel à sa vitesse. Si l'atome s'éloigne, les photons lui semblent, au contraire, de plus basse fréquence. Ce phénomène est nommé effet Doppler. De ce fait, à cause de l'agitation thermique, quelle que soit la précision avec laquelle on ajuste la fréquence du laser de couplage ou de l'impulsion lumineuse que l'on veut ralentir, tout se passe comme si le nuage de gaz recevait des faisceaux dont les fréquences sont mal ajustées.

Pour minimiser cette dispersion, nous utilisons des atomes extrêmement froids (qui se déplacent lentement). Quelques équipes avaient déjà obtenu



100

2. LE RALENTISSEMENT DE LA LUMIÈRE
éprouve sur la manipulation d'un nuage de
matière initialement opaque (un condensat
de Bose-Einstein) rendu transparent à l'aide
d'un faisceau laser.

les impulsions de lumière ralenties avant nous, mais elles utilisaient des gaz d'atomes à température ambiante et le ralentissement obtenu était limité. Nous avons créé les conditions nécessaires à l'obtention de la transparence électromagnétique induite dans un nuage d'atomes de sodium en forme de cigare – d'environ 0,2 millimètre de longueur et 0,05 millimètre de diamètre – piégé dans un champ magnétique et refroidi à moins d'un millionième de degré au-dessus du zéro absolu.

Nous refroidissons nos atomes de sodium à l'aide d'un dispositif combinant des faisceaux lasers, des champs magnétiques et des ondes radio. Le sodium est émis par un four chaud sous la forme d'un intense faisceau d'atomes propulsés à environ 2 600 kilomètres par heure. Ils sont alors heurtés de plein fouet par un faisceau laser qui, en une milliseconde, les ralentit à 160 kilomètres par heure – soit une décélération de 70 000 fois l'accélération de la pesanteur, g , produite par un faisceau laser qui ne brûlerait même pas les

doigts. La poursuite du refroidissement s'opère à l'intérieur d'une «mélasse optique», au moyen de six faisceaux qui baignent les atomes de tous côtés, et les ralentissent encore les atomes au point de les amener à 50 millionièmes de degré au-dessus du zéro absolu. En quelques secondes, 10 milliards d'atomes s'accumulent ainsi dans la mélasse optique. Nous éteignons alors les faisceaux laser, plongeant le laboratoire dans une obscurité totale, puis nous établissons des électroaimants dont le champ piège le nuage atomique. Durant 38 secondes, nous refroidissons les atomes par évaporation, éjectant les plus rapides d'entre eux pour ne garder que les plus lents, c'est-à-dire les plus froids. Des ondes radio minutieusement ajustées accélèrent l'éjection des atomes rapides. Tout ce phénomène – de la production du faisceau chaud jusqu'au confinement des atomes froids – se déroule à l'intérieur d'une chambre à vide où règne une pression 10^{14} fois inférieure à la pression atmosphérique (10 millionièmes de milliards d'atmosphères).

Une fois notre nuage refroidi à environ 500 milliardièmes de degré, il forme un condensat de Bose-Einstein, où les quelques millions d'atomes restants après l'évaporation se comportent de façon totalement synchrone. Le milieu le plus froid de l'Univers se trouve alors suspendu par des champs magnétiques au centre d'une chambre à vide.

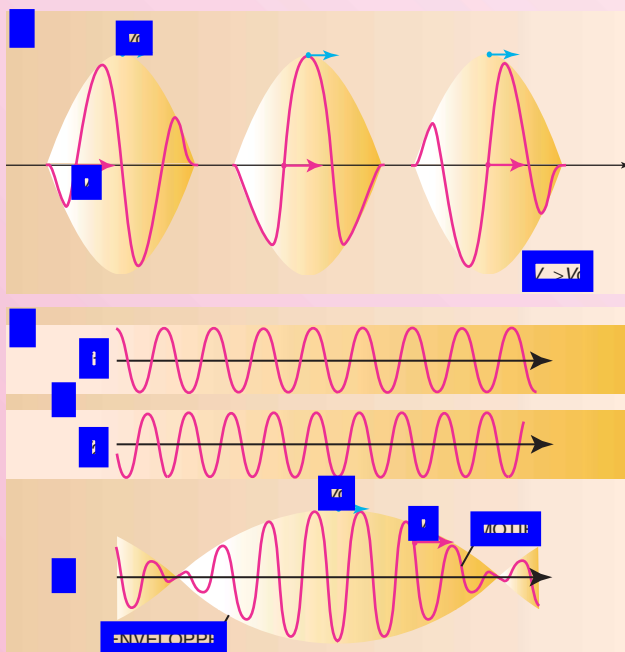
Le reste de notre dispositif expérimental, à un centimètre du nuage, est à la température ambiante. Des hublots hermétiques, sur les parois de la chambre, nous permettent d'observer les atomes en cours de refroidissement : un nuage atomique froid dans une mélasse optique ressemble à un petit soleil brillant, de cinq millimètres de diamètre.

Une fois notre cigare d'atomes en place, nous l'éclairons latéralement avec le laser de couplage, puis nous émettons une impulsion lumineuse le long son axe. Pour déterminer la vitesse de l'impulsion lumineuse dans le nuage, nous plaçons un détecteur de lumière derrière le nuage d'atomes et nous attendons que l'impulsion émerge pour

Vitesse de phase, vitesse de groupe

Une impulsion lumineuse est une oscillation électromagnétique de courte durée. Par définition la vitesse de phase est la célérité d'un point de cette oscillation. La vitesse de groupe, v_g , est la vitesse à laquelle progresse l'enveloppe tout entière de ce motif oscillant (voir la figure a).

Une impulsion lumineuse n'est jamais composée d'une seule fréquence, mais de la somme d'ondes sinusoïdales. L'intensité maximale de l'impulsion correspond au point où toutes ces ondes sinu-



soïdales sont en phase (où leurs crêtes coïncident). Quand l'indice de réfraction du milieu où se propagent ces ondes ne varie pas avec la fréquence, toutes les composantes «voyagent» à la même vitesse et le point où elles sont en phase également : la vitesse de groupe et la vitesse de phase sont égales. En revanche, quand l'indice de réfraction du milieu varie avec la fréquence, les composantes de l'impulsion se déplacent alors à des célérités différentes, de sorte que la vitesse de groupe et la vitesse de phase (v_p) sont différentes.

Dans le cas simple de deux ondes sinusoïdales de fréquences f_1 et f_2 , l'onde résultante est une oscillation comprise à l'intérieur d'une enveloppe dont la forme est une sinusoïde de période supérieure (voir la figure b). Cette onde résultante est le produit de deux ondes sinusoïdales que l'on désigne par l'onde «motif» et l'onde «enveloppe».

La fréquence de l'onde «motif» est la moyenne des fréquences f_1 et f_2 . Sa vitesse, qui n'est autre que la vitesse de phase, est donnée par l'indice de réfraction pour cette fréquence moyenne. Quant à l'onde «enveloppe», sa fréquence est proportionnelle à la différence entre f_1 et f_2 . On s'en convainc en se rappelant que si, à la limite, f_1 et f_2 sont égales, l'onde résultante est une sinusoïde «plate» (l'enveloppe est de fréquence nulle). En conséquence, la vitesse de groupe dépend également de l'écart entre les deux indices de réfraction $n(f_1)$ et $n(f_2)$, c'est-à-dire de la variation de l'indice de réfraction avec la fréquence.

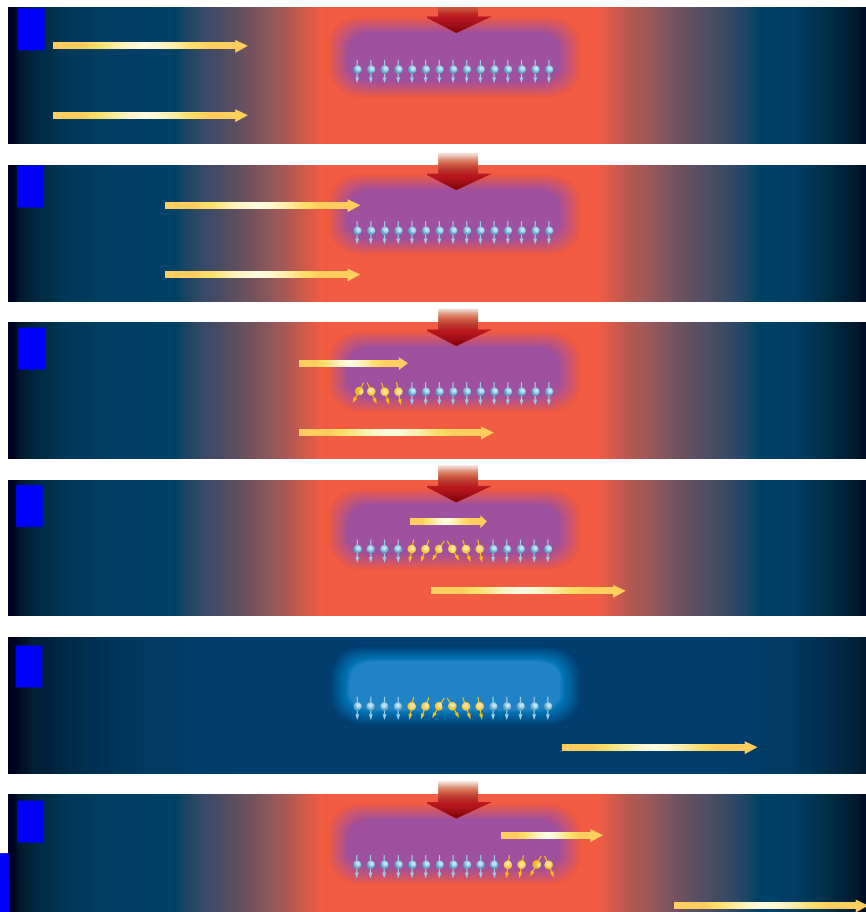
Ce raisonnement reste valable lorsque l'on ajoute de nombreuses composantes dont les fréquences sont, cette fois, séparées par des écarts infinitésimaux. La vitesse de phase est donnée par l'indice de réfraction à la fréquence moyenne de l'impulsion, et la vitesse de groupe dépend, de plus, de la variation de l'indice de réfraction. En fait, quand l'indice croît avec la fréquence, la vitesse de groupe est d'autant plus petite que l'indice varie rapidement.

mesurer le temps qu'il lui a fallu pour traverser. Immédiatement après, nous mesurons la longueur du nuage avec un autre faisceau laser qui l'éclaire par des sous et projette son ombre sur une caméra. Cette longueur divisée par le temps de traversée nous donne la vitesse de l'impulsion. Les durées de traversée obtenus varient de quelques microsecondes à quelques millisecondes. Pour la lumière, cela équivaut à un détournement dans un enroulement de plusieurs kilomètres de fibre optique.

Un gaz transparent

Comment le laser de couplage rend-il transparent le condensat d'atomes de sodium ? Il place les atomes du condensat dans un état quantique qui les rend incapables d'absorber la lumière de l'impulsion. Le sodium appartient à la famille des éléments alcalins qui regroupe les atomes n'ayant qu'un seul électron de valence (un seul électron « périphérique »). Cet électron de valence est responsable de la quasi-totalité des interactions du sodium avec son environnement. Lorsque l'électron de valence est sur son orbitale fondamentale, son énergie est minimale. Lorsqu'il quitte cette orbitale, il acquiert l'autant plus d'énergie qu'il s'éloigne du noyau : quand l'électron occupe une orbitale permise de grand rayon, l'atome est dans un état excité. Ces orbitales imposent la façon dont un atome de sodium interagit avec la lumière : c'est-à-dire les longueurs d'onde qu'il absorbe et qu'il émet.

Par ailleurs, l'électron de valence et le noyau de l'atome ont tous les deux un moment magnétique, c'est-à-dire qu'ils se comportent comme de minuscules aiguilles de boussole. De même que le champ magnétique terrestre est associé à sa rotation autour d'un axe Nord-Sud, le magnétisme de l'électron et celui du noyau sont associés à un moment angulaire intrinsèque, ou spin (une grandeur quantique dont l'équivalent, en physique classique, serait la rotation sur elles-mêmes des particules considérées). Lorsque les deux spins – les deux « aiguilles de boussole » – sont parallèles, mais de sens opposé (une configuration dite antiparallèle), l'interaction de leurs champs magnétiques est minimale, elle est maximale lorsqu'ils sont parallèles et de même sens. Par conséquent, outre le rayon de l'orbite occupée par l'électron de valence, l'alignement des spins



3. A L'INTERIEUR DU CONDENSAT : avant que l'impulsion n'atteigne le nuage d'atomes (en violet), tous les spins des atomes pointent dans la même direction (petites flèches). Un faisceau laser de couplage (en rouge) rend le nuage transparent à la lumière de même fréquence que l'impulsion (1, 2). Le nuage ralentit l'impulsion et la comprime (3), tandis que l'état des atomes change sur son passage, formant une onde de polarisation qui accompagne l'impulsion. L'ensemble est nommé polariton. Quand l'impulsion est tout entière dans le nuage (4), on éteint le laser de couplage (5). La partie lumineuse de l'impulsion disparaît mais le polariton laisse dans les atomes une empreinte qui contient toute l'information nécessaire pour la ressusciter. Un peu plus tard (6), le laser de couplage est à nouveau allumé : l'impulsion est régénérée et recommence à se propager.

électronique et nucléaire déterminent l'énergie des états excités.

Bien qu'un atome puisse occuper une multitude d'états excités, nous n'en avons utilisé que trois pour ralentir la lumière. Nous les avons nommés états 1, 2, et 3. Après sa préparation, le nuage de sodium ultrafroid contient des atomes qui se trouvent tous dans l'état 1, l'état fondamental, où l'énergie est minimale : l'électron de valence est sur sa plus basse orbitale, son spin est antiparallèle au spin du noyau et, en outre, le magnétisme total de l'atome est antiparallèle et de sens opposé au champ magnétique que nous utilisons pour maintenir le nuage en place. L'état 2 ressemble beaucoup à l'état 1, mais les spins de l'électron et du noyau sont parallèles et de même sens, ce qui augmente légèrement l'énergie de l'atome d'environ 0,002 électron-volt). L'état 3 est environ 300 000 fois plus énergé-

que l'état 2 et s'obtient en propulsant l'électron de valence sur une orbitale plus éloignée. En retombant de l'état 3 vers l'état 1 ou vers l'état 2, les atomes émettent une lumière dont la fréquence correspond à la couleur jaune caractéristique des lampes à vapeur de sodium utilisées pour les éclairages publics.

L'impulsion lumineuse que nous voulons ralentir est réglée sur une fréquence correspondant à l'énergie nécessaire pour passer de l'état 1 à l'état 3 (la lumière de cette impulsion est jaune). Si, sans autre précaution, on envoyait une telle impulsion sur le nuage, les atomes l'absorberaient totalement et passeraient de l'état 1 à l'état 3. Rapidement ils se désexciteraient et retomberaient dans l'état 1 en émettant la lumière jaune correspondante, mais ils le feraient au hasard dans le temps et dans l'espace. Le nuage raverait une lueur dif-

use et toute information sur la forme de l'impulsion originelle serait perdue.

Le faisceau laser de couplage, est réglé sur la fréquence correspondant à la transition entre les états 2 et 3. Lorsqu'ils sont dans l'état 1, les atomes ne peuvent l'absorber. Lorsque l'impulsion lumineuse, réglée sur la transition 1-3, arrive sur le nuage, son action conjuguée à celle du laser de couplage place les atomes dans une superposition quantique des états 1 et 2. Dans ces conditions, chaque atome est «simultanément» états dans l'état 1 et dans l'état 2. L'état 1, seul, absorberait la lumière de l'impulsion, et l'état 2 absorberait celle du faisceau de couplage. Dans chaque cas, les atomes se retrouveraient dans l'état 3, puis émettraient ensuite cette lumière au hasard. Toutefois, lorsque l'on tente de faire subir aux atomes ces deux transitions en même temps, elles s'annulent mutuellement selon un phénomène nommé interférence quantique. Le mélange d'états 1 et 2 superposés est désigné sous le nom «d'état sombre». Les atomes ne peuvent «être éclairés» par aucun des deux faisceaux (ils restent «dans le noir»). Le nuage devient transparent à la lumière de l'impulsion parce que les atomes qui le constituent placés dans cet état sombre, ne peuvent l'absorber. L'état superposé sombre est caractérisé par les proportions du mélange des états 1 et 2. Cette proportion dépend, en chaque point du nuage, du rapport des intensités reçues de la part de l'impulsion et du faisceau de couplage. Toutefois, une fois que le système se trouve dans un état sombre, il s'adapte pour rester sombre même lorsque les intensités relatives des deux faisceaux

varient : en un point donné du nuage lorsque l'intensité reçue de l'impulsion augmente, la proportion d'état 2 dans le «mélange» d'états augmente.

À cause d'une interférence quantique du même genre, l'indice de réfraction du nuage vaut exactement un – comme dans le vide – pour la lumière dont la fréquence correspond très précisément à la transition vers l'état 3. Toutefois, aux fréquences voisines, l'indice de réfraction ne vaut plus exactement un.

Le ralentissement de la lumière

Si l'indice de réfraction vaut un pour la fréquence correspondant à la transition vers l'état 3, notre impulsion réglée sur cette fréquence, ne devrait-elle pas progresser à la vitesse de la lumière dans le vide ? Non, parce que l'indice de réfraction du nuage varie très rapidement pour des fréquences voisines de cette fréquence fondamentale et qu'une impulsion lumineuse ne contient jamais une seule fréquence mais un intervalle de fréquences.

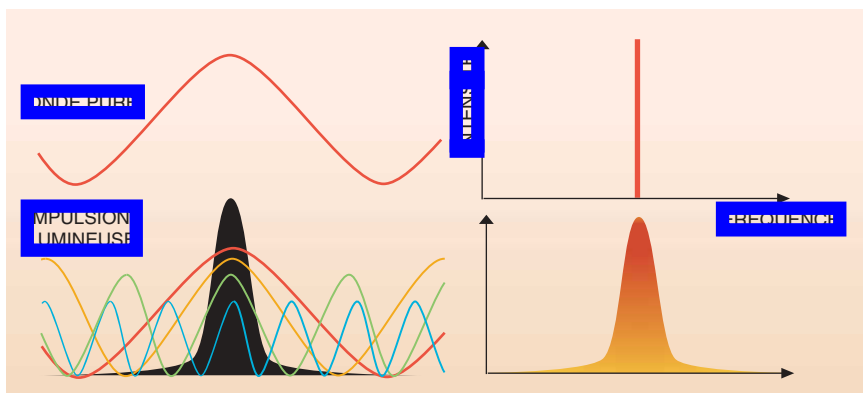
En effet, on peut considérer une impulsion comme une superposition d'oscillations sinusoïdales de fréquences voisines. Le pic d'intensité de l'impulsion correspond à l'endroit où toutes ces oscillations sont en phase où toutes leurs «crêtes» coïncident (voir la figure 4). La vitesse à laquelle progresse ce point, que l'on nomme «vitesse de groupe», n'est pas toujours égale à la célérité des ondes sinusoïdales – la vitesse de phase – constituant l'impulsion. Dans le vide, la vitesse de phase vaut environ 300 000 kilomètres par seconde, pour toutes les fréquences. Dans ce cas, toutes les ondes sinusoï-

dales constituant l'impulsion avancent à la même vitesse, et le point où elles sont en phase également. L'impulsion progresse alors aussi à la vitesse de la lumière dans le vide, vitesse de phase et vitesse de groupe sont égales. Toutefois, lorsque l'indice de réfraction varie avec la fréquence, les différentes ondes constituant l'impulsion se déplacent à des vitesses légèrement différentes. De ce fait, le déplacement du pic d'intensité de l'impulsion est modifié. En fait la vitesse de groupe, celle à laquelle progresse l'impulsion, diminue d'autant plus vite que l'indice de réfraction varie rapidement avec la fréquence (voir l'encadré page 72).

Cette façon de ralentir la lumière est très différente de ce qui se passe dans un milieu transparent ordinaire d'indice de réfraction supérieur à un. Tout d'abord, la vitesse de groupe est ralentie, mais la vitesse de phase est très peu modifiée, car l'indice de réfraction reste proche de un (dans un milieu transparent ordinaire, la vitesse de groupe diminue parce que toutes les vitesses de phase des ondes sinusoïdales diminuent). Dans notre cas, la vitesse de groupe diminue parce que l'indice varie brusquement sur un intervalle de fréquences étroit (voir la figure 5). Par ailleurs, le milieu ne conserve ces propriétés que si le laser de couplage est allumé.

Lorsque l'on diminue la vitesse d'une impulsion lumineuse d'un facteur égal à 20 millions, on observe plusieurs autres phénomènes. Avant d'entrer dans le condensat, l'impulsion lumineuse est longue de un kilomètre et se propage dans l'air à une vitesse proche de 300 000 kilomètres par seconde (bien sûr, notre laboratoire est loin d'avoir une longueur de un kilomètre, mais si nous plaçons notre laser à cette distance, ses impulsions auraient effectivement cette taille). Le front avant de l'impulsion traverse le hublot de verre, pénètre dans la chambre à vide puis dans le condensat d'atomes de sodium en suspension. À l'intérieur de ce minuscule nuage, la lumière se déplace à 54 kilomètres à l'heure (un coureur cycliste pourrait la dépasser).

Puisque le front avant de l'impulsion se déplace très lentement dans le nuage, alors que la partie arrière de l'onde arrive à toute allure, l'impulsion se tasse, telle un accordéon, dans le nuage de sodium. Sa longueur, comprimée d'un facteur 20 millions, n'est plus que de un vingtième de millimètre.



1. UNE IMPULSION LUMINEUSE est formée de plusieurs fréquences, contrairement à une onde sinusoïdale «pure». Le pic de l'impulsion correspond au point où toutes les ondes de fréquences différentes qui la composent sont en phase. Plus l'impulsion est localisée, plus son spectre est large. La vitesse à laquelle se déplace l'impulsion n'est égale à la célérité des ondes sinusoïdales que si elles se déplacent de façon synchrone, c'est-à-dire si l'indice de réfraction ne varie pas avec la fréquence.

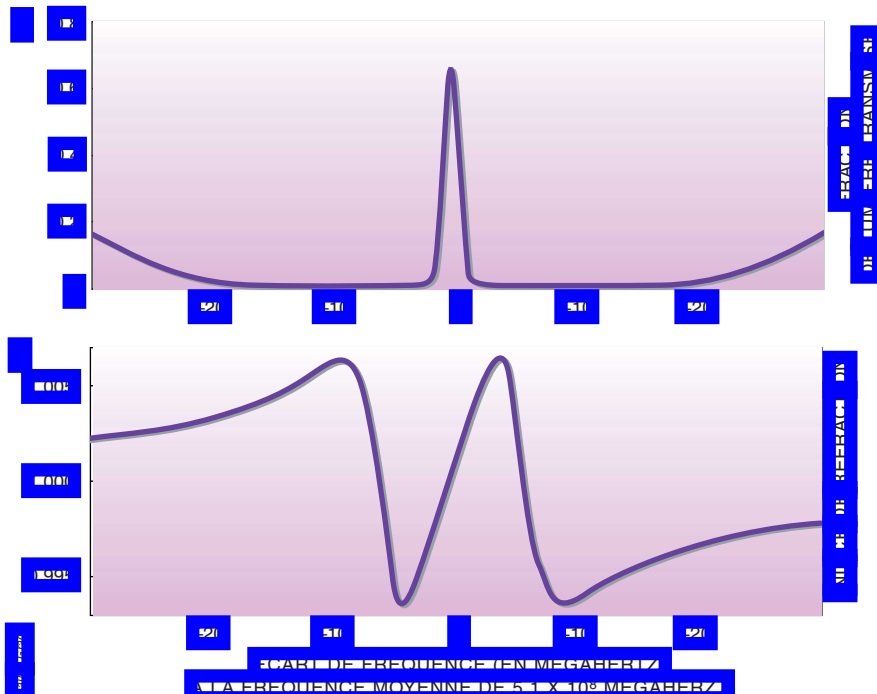
L'impulsion ainsi concentrée ne devrait-elle pas devenir extraordinairement intense ? Il n'en est rien : l'intensité lumineuse reste constante. Plus précisément, dans le vide, l'impulsion contient 50 000 photons ; dans le nuage elle ne contient plus sous forme d'énergie lumineuse, que l'équivalent de un quart de photon (encore le facteur 20 millions). Qu'est-il arrivé à l'énergie des autres photons ? Une partie de cette énergie a été dissipée dans le nuage et cédée aux atomes de sodium mais l'essentiel a été transféré au faisceau de couplage. Nous avons surveillé l'intensité du laser de couplage et nous avons observé ce transfert d'énergie.

L'empreinte de la lumière

À quoi ressemble l'impulsion lumineuse lorsqu'elle progresse dans le nuage ? Les transferts d'énergie entre les deux faisceaux lumineux modifient les états des atomes de sodium situés sur le passage de l'impulsion. Au niveau du front avant, les atomes passent de leur état originel – l'état 1 – à une superposition d'états 1 et 2 (l'état sombre). Dans la partie centrale de l'impulsion (la plus intense), la proportion d'états 2 contenue dans le « mélange » d'états est maximale. Après le passage de l'impulsion, les atomes reviennent dans l'état 1. Comme ces états correspondent à différents alignements des spins des atomes, c'est-à-dire à différentes polarisations, à chaque mélange d'états de proportion donnée correspond une polarisation particulière des atomes de sodium. Ainsi, l'impulsion de lumière lente comprimée dans le nuage apparaît accompagnée d'une « onde de polarisation » qui se propage dans le gaz. En fait, l'onde de polarisation et l'impulsion lumineuse se déplacent dans le condensat comme une seule et même entité nommée polariton.

Lorsque le polariton atteint l'extrémité du nuage de gaz, l'impulsion lumineuse se libère dans l'air et retrouve à la fois sa vitesse usuelle de 300 000 kilomètres par seconde, sa longueur initiale de un kilomètre et son énergie, qu'elle récupère dans le laser de couplage.

La vitesse de la lumière obtenue dans ces expériences dépend de plusieurs paramètres. Si certains sont fixés lorsque l'on a choisi l'espèce atomique et les états d'excitation utilisés, deux variables restent ajustables : la densité du nuage atomique et l'intensité du fais-



5. LES PROPRIÉTÉS OPTIQUES INDUITES dans le nuage de sodium par l'action conjuguée du laser de couplage et de l'impulsion ralentie permettent d'obtenir des facteurs de ralentissement de la lumière notables. Le nuage est rendu transparent pour la lumière d'une fréquence précise (a) et son indice de réfraction varie brusquement sur un court intervalle de fréquence (b). La transparence permet le passage sans absorption de l'impulsion lumineuse et sa vitesse est d'autant plus faible que la variation d'indice est brutale.

ceau laser de couplage. L'accroissement de la densité du nuage diminue la vitesse de la lumière, mais, dans les nuages très denses, les atomes sont plus difficiles à confiner dans un champ magnétique et s'échappent rapidement du piège. On peut également réduire la vitesse de l'impulsion en atténuant l'intensité du faisceau de couplage. Toutefois, là encore, le procédé est limité par le fait que si ce faisceau est trop faible, on perd l'effet de transparence électromagnétique induite ; le nuage devient opaque et absorbe l'impulsion.

On peut cependant obtenir le ralentissement ultime, l'arrêt total de la lumière, sans perdre l'impulsion par absorption : on éteint le faisceau de couplage pendant que l'impulsion ralentie et comprimée, se trouve toute entière dans le condensat, au cœur du nuage. L'impulsion lumineuse s'arrête et ne contient plus d'énergie lumineuse. Le polariton laisse sur les orientations des atomes du gaz froid une empreinte stationnaire qui constitue une sorte d'enregistrement de l'impulsion. Le polariton n'est pas déformé au cours de l'opération car, contrairement à l'impulsion lorsqu'elle pénètre dans le gaz, il ralentit d'un seul bloc. L'empreinte laissée sur les atomes contient toute l'information nécessaire pour reconstituer l'impulsion lumi-

neuse originelle. Ainsi, par exemple, l'intensité de l'impulsion en chaque point est donnée par les proportions des états 1 et 2. En somme, c'est un hologramme de l'impulsion qui se trouve imprimé dans le condensat.

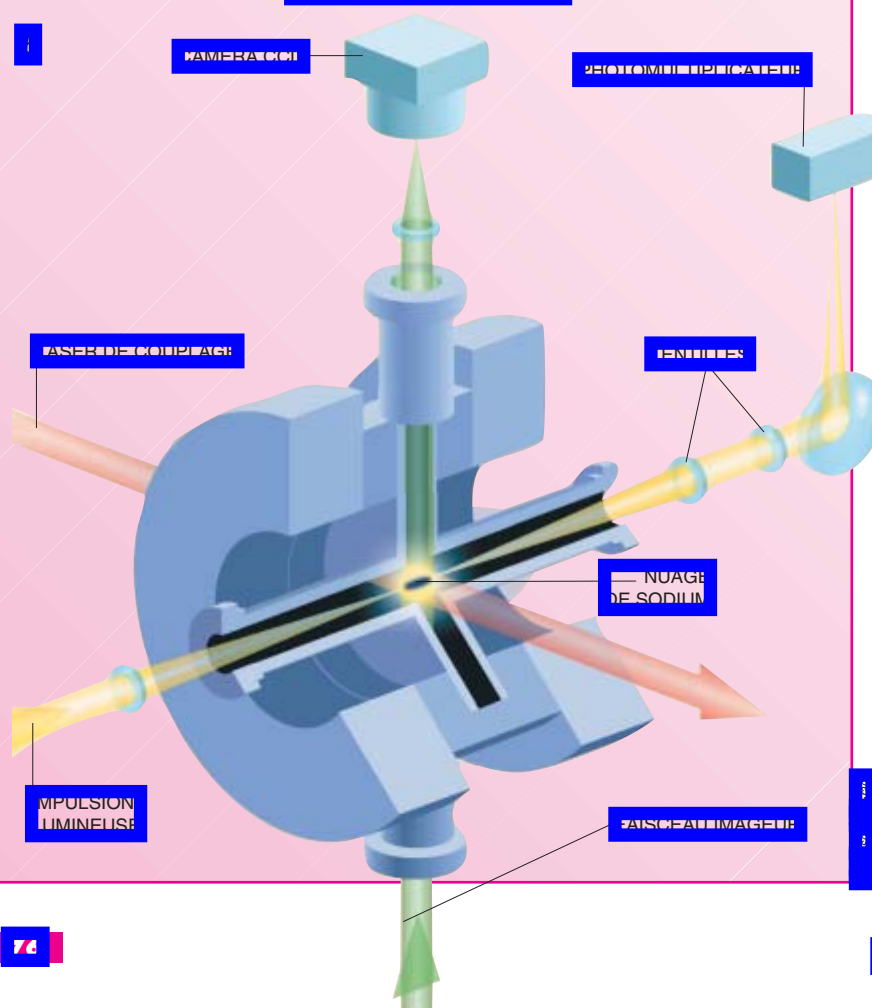
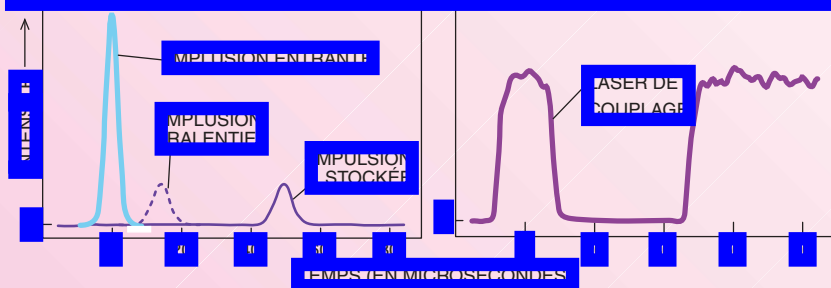
Lorsqu'on l'éclaire à nouveau le condensat avec le laser de couplage, l'impulsion lumineuse réapparaît et se propage à nouveau comme si rien ne l'avait interrompue. Toutefois, on ne peut stocker la lumière que durant une milliseconde tout au plus (si elle n'avait pas été piégée, elle aurait parcouru pendant ce temps 300 kilomètres dans l'air). Plus l'impulsion reste stockée longtemps, plus elle se dégrade : les atomes du gaz restent, même dans le condensat, animés de mouvements erratiques qui finissent par brouiller la configuration de polarisation. En outre, les collisions entre atomes peuvent dégrader les états de superposition. Au bout d'une milliseconde, l'impulsion de sortie est bien plus faible que l'impulsion originelle, mais nous avons trouvé un moyen d'y remédier. Si, lorsque l'on rallume le faisceau de couplage, on lui confère une intensité supérieure, l'impulsion de sortie est plus brillante, mais plus courte. Si l'on active et désactive rapidement le faisceau de couplage à plusieurs reprises, on régénère l'impulsion en plusieurs « morceaux ». De

Stopper la lumière

L'expérience de ralentissement de la lumière repose sur trois faisceaux laser et un nuage d'atomes de sodium ultrafroids (exagérément grossi sur le schéma) enfermé dans une enceinte où règne un vide intense. Le faisceau de couplage interagit avec le nuage et le rend transparent à la lumière du laser qui produit les impulsions. Rendu transparent, ce nuage ne ralentit pas la lumière que très lentement.

Un tube photomultiplicateur mesure le temps d'arrivée de l'impulsion. Le faisceau d'imagerie mesure alors la longueur du nuage en projetant son ombre sur une caméra. Plusieurs éléments ne sont pas représentés sur le schéma : le dispositif de création et de refroidissement d'un nouveau nuage pour chaque impulsion ; les électroaimants dont les champs combinés maintiennent les atomes en place et certains dispositifs optiques annexes.

On observe le ralentissement des impulsions lumineuses en comparant de façon très précise les instants où elles sont détectées. L'instant où l'on détecte l'impulsion en l'absence de condensat correspond à l'instant zéro. Le ralentissement des impulsions par la présence du condensat est mis en évidence par le retard avec lequel elles atteignent le détecteur (*en pointillés sur la courbe de gauche*). Pour stopper complètement une impulsion, on éteint le faisceau de couplage au moment où elle se trouve à l'intérieur du nuage. Le temps d'arrêt de l'impulsion – environ 35 microsecondes – s'ajoute à son retard. Une impulsion ralentie sort du dispositif avec une intensité diminuée car le nuage n'est pas parfaitement transparent. En outre, une impulsion stoppée se dégrade petit à petit sous l'action de diffusions et de collisions parmi les atomes qui la retiennent.



elles manipulations donnent une idée de la précision avec laquelle on contrôle les impulsions stockées.

Trous noirs et ordinateurs

Le ralentissement et l'arrêt de la lumière ouvrent la voie à de nombreuses expériences intéressantes. On pourrait, par exemple, injecter, dans un condensat de Bose-Einstein, une impulsion lumineuse dont la vitesse égalerait celle du son dans ce condensat (environ un centimètre par seconde). Dans ces conditions, une onde sonore comparable au « bang » des avions accompagnerait l'impulsion et le condensat tout entier se mettrait à osciller. Ce serait une façon nouvelle d'étudier les propriétés superfluides des condensats (certains liquides sont superfluides à très basse température, c'est-à-dire que leur viscosité s'annule, de la même façon que la résistance s'annule dans les métaux très froids).

Aujourd'hui, on sait également créer des vortex (des tourbillons) dans un condensat en rotation. Une impulsion de lumière lente traversant un tel vortex serait entraînée par le mouvement du gaz, un peu à la façon dont les rayons lumineux sont déviés au voisinage des astres massifs ou des trous noirs en rotation. Grâce à la lumière lente, on pourrait simuler ces phénomènes en laboratoire.

La lumière lente annonce également une nouvelle forme d'optique non linéaire fondée sur l'interaction de deux faisceaux laser. Beaucoup de phénomènes d'optique non linéaire sont produits lorsque l'indice de réfraction du milieu considéré varie avec l'intensité de la lumière reçue. Cela ouvre un vaste champ de recherches, tant fondamental qu'appliqué, à des domaines aussi variés que l'imagerie ou les télécommunications. L'obtention de certains phénomènes d'optique non linéaire exige aujourd'hui des faisceaux extrêmement intenses, mais, grâce à la lumière ralentie, on pourra les produire avec un très petit nombre de photons. Ces phénomènes devraient être utiles pour la construction d'aiguillages ultrasensibles destinés aux futurs réseaux Internet entièrement optiques.

La lumière ralentie et piégée pourrait également être utilisée pour la construction d'ordinateurs quantiques. Les ordinateurs quantiques seraient des calculateurs où les bits habituels de l'informatique, qui valent soit 0 soit 1

seraient remplacés par des qubits (pour *quantum bits*), des superpositions de 1 et de 0. Un système de n qubits exploiterait «simultanément» 2^n états différents et serait utilisé pour réaliser un très grand nombre d'opérations en parallèle. Ces ordinateurs, s'ils sont un jour opérationnels, résoudront en peu de temps des problèmes mathématiques qui nécessiteraient des durées de calcul considérables sur les ordinateurs ordinaires. Dans les systèmes que l'on envisage aujourd'hui, il existe deux grandes familles de qubits. Les premiers sont fixés dans l'espace : ce sont par exemple des groupes d'atomes dont les spins nucléaires pointent dans différentes directions et qui interagissent rapidement les uns avec les autres. Les seconds se déplacent rapidement d'un endroit à un autre (ce sont, par exemple des photons), mais sont difficiles à faire interagir selon les modes exigés par un ordinateur quantique. En transformant des photons mobiles en polaritons stationnaires – et inversement –, les dispositifs à lumière piégée fourniraient un moyen fiable d'opérer une conversion mutuelle entre ces deux types de qubits. On pourrait par exemple, imprimer deux impulsions au sein d'un même nuage d'atomes : cette impression déclencherait l'interaction de polaritons, c'est-à-dire, d'un point de vue informatique, le traitement de l'information contenue dans les impulsions d'entrée. La lecture du résultat s'obtiendrait en transformant l'empreinte laissée par les polaritons en de nouvelles impulsions lumineuses de sortie.

Même si la lumière ralentie n'est pas le moyen le plus pratique ni le plus universel pour la construction de l'ordinateur quantique, elle aura ouvert de nouvelles voies de recherches.

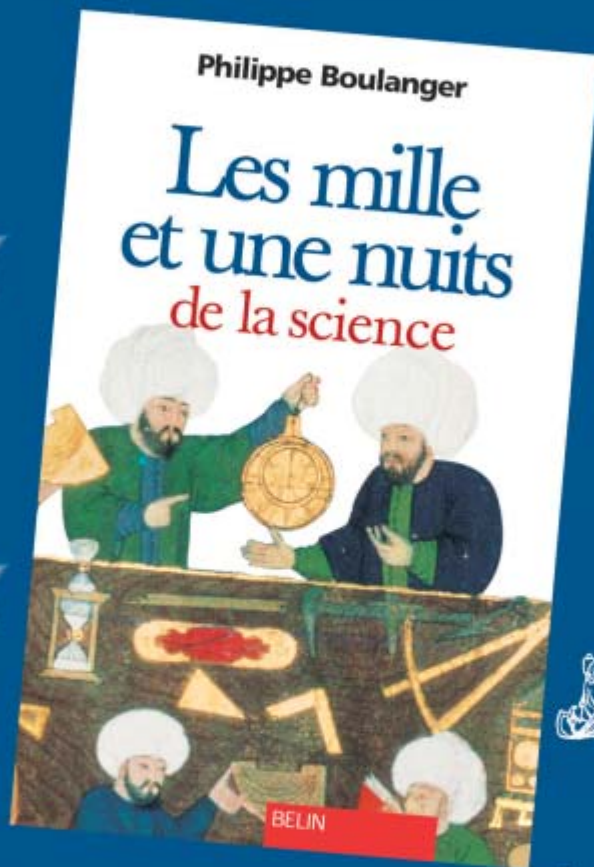
Lene Versteegard HAU, professeur de physique à l'Université de Harvard dirige le groupe Atomes refroidis de l'Institut Rowland pour les sciences de Cambridge, dans le Massachusetts

Stephen E. HARRIS, *Electromagnetically Induced Transparency*, in *Physics Today* vol. 50, n° 7, pp. 36-42, juillet 1997.

Eric A. CORNELL et Carl E. WIEMAN, *La condensation de Bose-Einstein*, in *Pour la Science*, n° 247, mai 1998.

Chien LIU, Zachary DUTTON, Cyrus H. BEHROOZI et Lene Versteegard HAU, *Observation of Coherent Optical Information Storage in an Atomic Medium Using Halted Light Pulses*, in *Nature* vol. 409, pp. 490-493, 25 janvier 2001.

**Entrez dans l'ambiance orientale
des mille et une nuits et laissez-vous
charmer par l'imagination
de Shéhérazade contant la science
et les tribulations des savants**



«Dans ce livre, Philippe Boulanger, Directeur de la revue Pour la Science, réussit le tour de force de vulgariser des concepts difficiles en charmant son lecteur par une subtile danse... des neurones».

Hervé Ponchelet
Le Point

Code 002445 • 240 pages • Prix 100 F
Voir bon de commande p. 96



BELIN • POUR LA SCIENCE

Or et céramiques des Andes du Nord

JEAN-FRANÇOIS BOUCHARD

Les pillards ont longtemps «exploité» la nécropole équatorienne de Tumaco-La Tolita, qui regorgeait de figurines et de bijoux en or. Depuis quelques années, les archéologues ont enfin accès à cette région marécageuse et reconstituent l'histoire des Amérindiens qui y ont prospéré vers le début de notre ère.

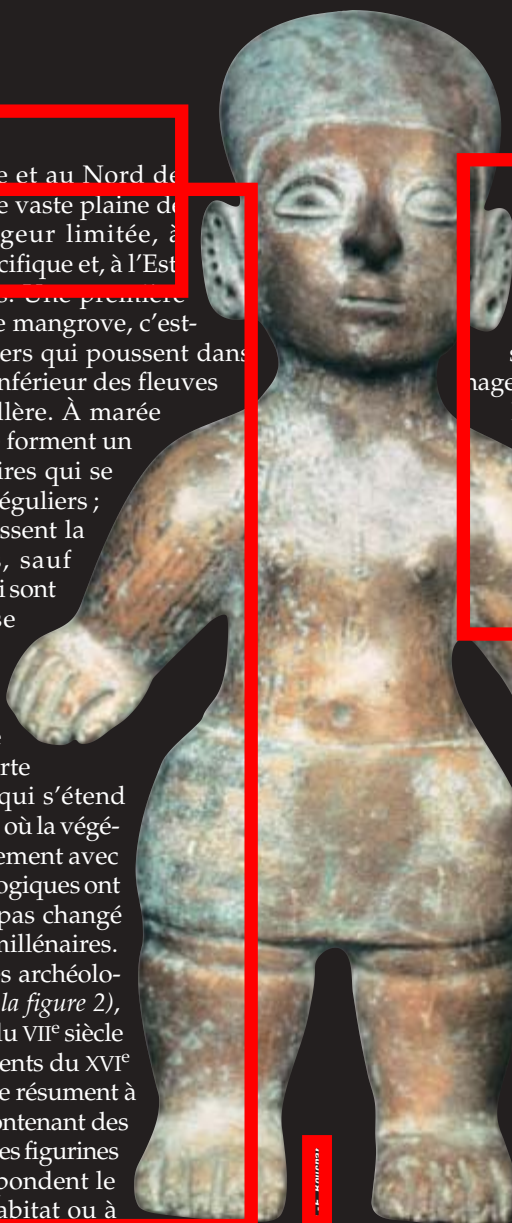
Au Sud de la Colombie et au Nord de l'Équateur, s'étend une vaste plaine de 50 kilomètres de largeur limitée, à l'Ouest, par l'océan Pacifique et, à l'Est, par la cordillère des Andes. Une première frange côtière, recouverte de mangrove, c'est-à-dire de forêts de palétuviers qui poussent dans les vasières, borde le cours inférieur des fleuves qui descendent de la cordillère. À marée haute, les chenaux de marée forment un dédale d'affluents temporaires qui se jettent dans les cours d'eau réguliers ; tout autour, les sols où poussent la mangrove sont immergés, sauf quelques-uns, plus élevés, qui sont les seuls où l'on puisse construire des habitations. Au-delà de cette frange côtière de plusieurs kilomètres de largeur commence une plaine alluviale recouverte de forêt tropicale humide, qui s'étend jusqu'au pied de la cordillère, où la végétation se modifie progressivement avec l'altitude. Les études palynologiques ont montré que ce paysage n'a pas changé au cours des trois derniers millénaires.

La région regorge de sites archéologiques préhispaniques (voir la figure 2), dont les plus anciens datent du VII^e siècle avant notre ère et les plus récents du XVI^e siècle de notre ère. Ces sites se résument à des couches sédimentaires contenant des vestiges archéologiques, tels des figurines ou des récipients. Ils correspondent le plus souvent à des sols d'habitat ou à

des dépotoirs anciens. Ces derniers contiennent plus de 95 pour cent de fragments de céramiques, car, dans ce milieu humide, les vestiges organiques ne se sont pas conservés. On a également retrouvé des sites agricoles (voir la figure 3), caractérisés par des talus surélevés, séparés par des canaux de drainage. La présence de tessons de céramique dans les sédiments suggère que les populations amérindiennes y résidaient aussi, au moins de façon épisodique. Enfin, le site le plus important est la nécropole de l'île de La Tolita, dont les centaines de tombes contenaient d'innombrables objets en or (voir la figure 1), un site aujourd'hui dévasté par les fouilles

clandestines

Qui a façonné ces objets ? Comment les habitants de ce environnement si particulier où l'eau est omniprésente, s'y sont-ils adaptés ? Ils appartenaient à la culture dite de «Tumaco-La Tolita» qui tire son nom du site archéologique de l'île de Tumaco, au Sud de la Colombie, et de la nécropole de l'île de La Tolita, dans l'estuaire du río Santiago, sur la côte Nord de l'Équateur. Elle s'étendait de Buenaventura à Esmeraldas. Cette culture a atteint son apogée entre 300 avant notre ère et 300 de notre ère. Les Indiens qui y appartenaient ont utilisé l'océan et les cours d'eau du littoral à la fois pour se nourrir, complétant par la pêche leurs ressources alimentaires issues de l'agriculture et de la cueillette, et pour se déplacer beaucoup plus aisément.



que par voie terrestre, ce qui leur évitait de traverser des zones de forêt dense et de marécages. Ils maîtrisaient les techniques de navigation. Leur société devait être organisée autour de rites funéraires et de la collecte de l'or dans la plaine alluviale. Les objets mis au jour témoignent de la splendeur de ces populations indigènes qui ont prospéré dans les Andes du Nord pendant plusieurs siècles.

Il y a 25 ans, on savait peu de choses du passé précolombien de la région, hormis que ses habitants avaient produit des objets en or et en céramique qui parvenaient dans les collections publiques et privées par l'intermédiaire des pillards. Au milieu du XX^e siècle, sur l'île de La Tolita, les excavations ont même atteint un rendement quasi industriel : le propriétaire de l'île y avait installé un système de lavage de la terre pour récupérer les objets ou les fragments d'objets en or qui s'y trouvaient, comme s'il s'agissait d'un gisement alluvial. Immediatement les conquistadors du XVI^e siècle, qui fondaient les idoles et les bijoux des peuples amérindiens pour enrichir les rois d'Espagne, lui auraient envoyé plusieurs tonnes d'or en Italie sous forme de lingots. Outre que les sites étaient farouchement gardés par les pillards, la géographie et le climat de la région, peu propices à la fouille, n'incitaient guère les archéologues à s'y rendre. De plus, on connaissait tout aussi mal les autres cultures des Andes du Nord, de sorte que l'on ne savait pas situer la culture de Tumaco-La Tolita dans un cadre historique.

Une terre vierge de toute archéologie

En 1976, une mission archéologique française s'est établie en Colombie puis en Équateur. Nous avons alors prospecté dans la région, pour repérer les signes émus de la présence d'un site archéologique. Seul un œil exercé est alerté par un monticule artificiel, unesson affleurant à la surface... À ces endroits, nous avons effectué des sondages, pour confirmer la présence de vestiges archéologiques, avant d'entreprendre des fouilles. Nous avons également fouillé des sites urbains, découverts lors de chantiers de génie civil, notamment à Tumaco. Enfin, avec nos collègues équatoriens, nous avons travaillé à différents endroits de l'île de La Tolita, à l'écart des sépultures précolombiennes saccagées par les pillards. Dans ces sols gorgés d'eau, nous avons pompé l'eau d'infiltration qui remplissait la fouille à mesure que nous progressions. Malgré ces dispositions, nous avons été contraints de travailler plus rapidement qu'à l'accoutumée. Nous avons poursuivi les fouilles jusqu'en 2000, quand la mission française a pris fin.

Par ailleurs, nos collègues équatoriens, colombiens, espagnols et nord-américains ont fouillé d'autres sites archéologiques : de petits hameaux ou des villages et, plus rarement des agglomérations de grande taille, tel le site d'Atacames près d'Esmeraldas.

LE STYLE DE CES FIGURINES est typique de la culture de Tumaco-La Tolita. La figurine en céramique (à gauche) a été façonnée entre le III^e siècle avant notre ère et le III^e siècle de notre ère. Le personnage est habillé d'un pagne court ; la forme fuyante de sa coiffure évoque une déformation crânienne (qui se pratiquait dans d'autres cultures amérindiennes, notamment chez les Mayas). La figurine en or (à droite) a un œil en platine, comme d'autres objets de la culture de Tumaco-La Tolita, qui allient les deux métaux précieux. Elle ne provient pas de la zone côtière, comme la précédente mais du site d'El Angel, situé dans la sierra andine.



Aujourd'hui, d'après le travail des archéologues des différentes missions, réalisé pendant les 25 dernières années, on peut désormais ébaucher quelques caractéristiques de la culture de Tumaco-La Tolita.

A partir du style des fragments de céramiques retrouvés dans divers sites et du contexte stratigraphique de leur découverte, nous avons précisé la chronologie des collections d'objets d'archives publiques et privées. La datation au carbone 14 de charbons présents dans les sédiments a encore affiné ces résultats.

La première phase d'occupation se situe entre 700 et 500 avant notre ère environ. Elle correspond au peuplement de la région par des Indiens venus des côtes équatoriennes plus méridionales comme en atteste la ressemblance stylistique des céramiques avec celles de

Entre 500 et 300 avant notre ère, La Tolita devient une importante métropole qui exerce son hégémonie sur un vaste territoire côtier de plus de 500 kilomètres de longueur. Les objets en céramique, tels les récipients à usage domestique ou cérémoniel, les figurines à forme humaine et animale (ou en forme de chimères, à mi-chemin entre l'homme et l'animal) sont abondantes. Leur style, notamment celui des figurines, est homogène dans toute la région située entre Buenaventura et Esmeraldas Atacames. On a également retrouvé des figurines du même style en dehors de ces limites, par exemple dans la Sierra du Nord, près de Quito, ou sur le piémont Ouest (la pente douce au pied de la cordillère où les matériaux que l'érosion a arrachés à la montagne se sont accumulés) : l'influence de la culture de Tumaco-La Tolita est probablement liée à un fort pouvoir politique et religieux.

À La Tolita, cette culture a laissé d'innombrables bijoux (*voir la figure 6*) et d'autres ornements en or dans les sépultures de la nécropole. Cette orfèvrerie est l'une des plus anciennes des Andes du Nord : les autres cultures de la région ont plutôt prospéré à partir du 1^{er} siècle de notre ère (en Colombie). Par ailleurs, la culture de Chorrera n'a pas laissé d'orfèvrerie. La société de Tumaco-La Tolita devait être structurée autour de rites funéraires et de la collecte de l'or dans toute la plaine alluviale. Ce territoire de mangrove et de forêt tropicale n'a sans doute été occupé que parce qu'il y avait beaucoup du précieux métal.

Entre l'océan et la cordillère, le territoire habité était relativement restreint. La plupart des sites d'habitat ont été découverts dans une étroite frange de la plaine alluviale, entre la mangrove et la forêt, principalement le long des cours d'eau de la partie aval de la plaine. Les Amérindiens de la culture de Tumaco-La Tolita ont probablement fait des incursions en haute mer ou dans le piémont et les bas versants qui entouraient son territoire pour pêcher ou se procurer des matériaux destinés à la fabrication des objets de la vie domestique. Toutefois, la plupart du temps, ils devaient se limiter à la mer proche des côtes, aux estuaires et aux cours d'eau, à la mangrove et à la forêt de la plaine alluviale intérieure pour y chasser et y pêcher.



2. L'aire occupée par les populations de la culture de Tumaco-La Tolita se trouvait entre ces deux sites éponymes, où se concentraient la majorité des sites archéologiques. Toutefois, son influence, dont témoignent les objets découverts en dehors de ces limites, s'étendait jusqu'à Buenaventura, au Nord, et Esmeraldas, au Sud. Elle était installée sur la frange côtière recouverte de mangrove et de forêt tropicale, entre l'océan Pacifique à l'Ouest et la cordillère des Andes à l'Est.

Les Andes équatoriales à l'époque préhispanique

En Amérique du Sud, l'Équateur est l'un des premiers endroits où les hommes préhistoriques ont adopté l'agriculture à la place des modes de subsistance traditionnels, la chasse, la pêche et la cueillette. Ce processus, que l'on nomme la néolithisation, s'est déroulé au cours des quatrième et troisième millénaires avant notre ère, sur la côte Pacifique.

On attribue l'avènement du Néolithique aux Amérindiens de la culture Valdivia, qui ont prospéré en Équateur pendant un peu plus de deux millénaires, de 3800 à 1500 avant notre ère environ. Les Amérindiens de Valdivia cultivaient le maïs, le haricot, la courge, la Calebasse et le coton. En outre, ils ont inventé la céramique et le tissage. Les sites archéologiques montrent que, sur leurs lieux d'habitation et aux alentours, l'espace était organisé, et que des relations hiérarchisées existaient entre les centres politiques cérémoniels et les villages qui les entouraient. Le site de Real Alto, le plus important que ces groupes amérindiens de Valdivia aient laissé, à partir duquel on a défini les phases successives du développement de la culture Valdivia, en fournit un exemple. La première occupation correspond à un hameau de quelques huttes, dont les habitants vivaient de la chasse, de la pêche et de la cueillette, mais aussi de la culture du maïs, des légumineuses comme le haricot et le pois, des bulbes et des racines comestibles, tel le manioc. Puis ils commencent à pratiquer l'agriculture sur brûlis, construisent de grandes maisons de plan elliptique disposées sur trois côtés d'un espace rectangulaire. La culture du maïs s'intensifie. La taille des champs augmente, ainsi que leur éloignement par rapport à Real Alto. Un réseau de villages satellites apparaît, qui dépendent de ce centre, devenu la résidence d'une élite politique et religieuse. De telles structures socio-économiques politiques et rituelles sont les premières des Andes équatoriales. Enfin, les techniques de culture s'améliorent : les Amérindiens des dernières phases de Valdivia pratiquent la culture sur talus surélevés, la culture en terrasses, et construisent des retenues d'eau. Des formes améliorées des plantes cultivées apparaissent. Au total, les rendements augmentent.

L'art plastique de la culture Valdivia est le plus ancien du continent américain. Les artisans façonnent des figurines anthropomorphes en pierre, puis en argile. Elles sont formées de deux bouddins modelés et accolés et représentent surtout des femmes, parfois enceintes et parfois bicéphales.

Entre 1450 et 1350 avant notre ère, la culture Machalilla, dont on connaît peu de choses, succède à la culture Valdivia. Puis, entre 1000 et 850 avant notre ère, apparaît la culture Chorrera, qui unifie les groupes indigènes de la région. Ils s'organisent en chefferies, la base de toutes les structures sociopolitiques futures dans les Andes équatoriales. La technique des potiers Chorrera surpasse celle de leurs ancêtres, ainsi que leur niveau artistique. Ils façonnent des récipients fonctionnels, mais aussi des vases en forme d'animaux et de plantes, ainsi que des statuettes représentant des personnages debout ou assis en tailleur, à l'attitude hiératique.

Vers 500 avant notre ère, la culture Chorrera donne naissance à de nombreuses cultures régionales, dont chacune se caractérise par un style artistique propre. Les groupes d'Amérindiens correspondants ont les mêmes racines, mais forment probablement des entités ethniques et socio-politiques distinctes. La culture Tumaco, la Tolita, installée sur le littoral Pacifique du Nord de l'Équateur et du Sud de la Colombie, est l'une d'elles. Plus au Sud, on rencontre les cultures Jama-Coaque, Bahia de Caraquez, Guangala et Jambelí. D'autres encore fleurissent dans la sierra andine, y compris sur le



LA CULTURE VALDIVIA, qui a prospéré en Équateur pendant les quatrième, troisième et deuxième millénaires avant notre ère, est la première des Andes équatoriales qui ait adopté l'agriculture, inventé la céramique, façonné des objets d'art, telles ces représentations féminines modelées dans l'argile et cuites (qui mesurent environ dix centimètres de hauteur en haut), et des récipients à décoration anthropomorphe (en bas).

versant oriental des Andes, du côté de l'Amazonie. Toutes prospèrent jusqu'à environ 500 de notre ère.

Du VI^e au XV^e siècle, de grandes chefferies, probablement plus complexes que les précédentes, se constituent et établissent entre elles des relations économiques et politiques. Les groupes Manteños et Huancavilcas forment des « ligues maritimes » : ils commercent le long des côtes, avec de grands radeaux à voile (dont s'inspirera le zoologue norvégien Thor Heyerdal pour construire le *Kon Tiki*, sur lequel il traversera le Pacifique avec six autres aventuriers, en 1947), tandis que les groupes Canaris, Caranquis ou Caras s'épanouissent dans les territoires montagneux de la sierra.

À la fin du XV^e siècle, les Incas arrivent du Pérou dans les Andes équatoriales. Leur immense empire s'étend à toutes les Andes centrales. Ils conquièrent une partie des Andes du Sud et des Andes équatoriales. En Équateur, ils prennent le contrôle des terres agricoles et du commerce du spondyle, un coquillage coloré qui, à l'époque, était aussi précieux que les gemmes dans le monde occidental. Cependant, les Incas ne réussissent pas à établir leur suprématie sur la côte Nord, recouverte de marécages où les populations indigènes restent maîtresses de leurs territoires. Au XVI^e siècle, l'arrivée des conquistadors, marque la fin des cultures amérindiennes.



3. LES SITES ARCHEOLOGIQUES se résument souvent à un monticule artificiel (*ci-contre*), que les archéologues sondent pour confirmer la présence de vestiges du passé avant d'entreprendre des fouilles (la culture de Tumaco-La Tolita n'a laissé aucun monument). Les sites archéologiques agricoles (*ci-dessus*) se caractérisent par les canaux de drainage (les zones envahies par les herbes hautes) et des talus surélevés (qui forment des bandes parallèles aux canaux), où l'on pratiquait la culture du maïs et peut-être aussi celle du manioc.

Ils pratiquaient aussi l'agriculture. Dans les basses terres de la plaine alluviale, des aires naturellement humides et saturées par les pluies et les crues étaient drainées par des réseaux de canaux, ce qui permettait de cultiver le manioc sur des talus surélevés (voir la figure 3). Ces sites ne permettent pas d'estimer l'effectif des groupes de la culture de Tumaco-La Tolita. On ignore combien il y avait de sites et d'habitants par site au total. Cependant, le rendement était assez élevé pour que le travail d'une partie de la population suffise à produire de la nourriture pour la totalité du groupe, notamment pour diverses classes d'artisans, tels les potiers et les orfèvres, qui ont laissé les vestiges par lesquels nous connaissons cette culture. Ce n'est qu'en exploitant les différents milieux qui les entouraient que les représentants de la culture de Tumaco-La Tolita ont réussi à se développer. Par cette «extension à l'horizontale», qui multiplie les milieux exploités à partir d'un foyer d'habitat préférentiel, elle se différencie des autres cultures américaines de forêt tropicale et se rapproche de celles des régions montagneuses des Andes, dites d'économie verticale, qui essaient de contrôler les paliers écologiques situés à des altitudes diverses.

A partir de 300, la culture de Tumaco-La Tolita décline et, pendant tout le millénaire qui précède la conquête espa-

gnole, diverses cultures moins importantes et encore mal connues lui succèdent. Les quelques sites archéologiques retrouvés ne témoignent que d'une occupation intermittente de la région. La céramique devient plus rare, moins élaborée et l'orfèvrerie disparaît. La région est occupée par des groupes indigènes qui ont oublié la plupart des acquis de la culture Tumaco-La Tolita. À leur arrivée, les conquistadors découvrent des traditions culturelles et des techniques très en deçà de celles qui avaient auparavant existé dans ces régions, ce qui conforte leur opinion sur le pays : c'est un «enfer vert» peu propice au développement économique. Entre 1524 et 1527, lors de ses premières expéditions vers le Pérou à partir de Panama, le futur conquérant de l'Empire inca, Francisco Pizarre, s'arrête à l'île du Coq, près de Tumaco. Cependant, il ne tente pas de conquérir ces terres trop hostiles et passe ainsi à côté du fabuleux trésor enfoui dans les sépultures de La Tolita. Par la suite, les marécages, les forêts tropicales et un climat éprouvant, chaud et humide, sont autant d'obstacles à la pénétration des fan- tassins et des cavaliers espagnols, dont l'équipement est plus une gêne qu'une protection. Ils sont décimés par les maladies, les accidents et la résistance armée des

indigènes. En outre, les habitants de Buenaventura, au Nord et de Guayaquil, au Sud, n'ont pas intérêt à voir se développer un nouveau port, qui leur ferait de la concurrence : cette côte, déjà délaissée par les Incas lorsqu'ils avaient envahi l'Équateur, l'est aussi des conquistadores et des colons espagnols.

En revanche, dès les débuts de la colonisation espagnole, un premier groupe d'esclaves noirs parvient à s'échapper lors d'un naufrage et s'établit sur la côte d'Esmeraldas. Ensuite, de nombreux autres esclaves en fuite trouvent asile dans ces terres difficiles d'accès. Entre 1850 et 1920, les travailleurs des exploitations aurifères parcourent les alluvions des fleuves et ceux des exploitations forestières viennent grossir cette population, qui constitue désormais un important groupe afro-américain. Aujourd'hui, leurs descendants cohabitent avec les indiens, de moins en moins nombreux. Ces derniers, durement touchés par les maladies d'origine européenne au début de la conquête espagnole, ont perdu leurs anciens territoires et ont été repoussés dans les parties amont des cours d'eau, très hostiles, où les chances de prospérer sont encore plus réduites que sur la partie littorale de la plaine alluviale.

Navigation préhispanique

A l'intérieur de la région du littoral Nord-équatorial, on ne peut se déplacer et transporter des marchandises que par voie d'eau ; les rivières et les marécages sont autant d'obstacles à la progression terrestre. D'après les récits de conquistadores et les vestiges archéologiques, les embarcations préhispaniques étaient soit des radeaux, soit des pirogues. Dans la région, les conquistadores ne mentionnent que des pirogues de tailles diverses, faites d'une seule pièce de bois, ou de précaires plates-formes flottantes qui servaient d'embarcations temporaires en eaux calmes, mais pas de grands radeaux à voile destinés à naviguer en mer, tels qu'ils en avaient vus plus au Sud, le long des côtes du



Pérou. En dehors de ces témoignages du XVI^e siècle, nous ne disposons que de quelques figurines en céramique de la culture Tumaco-La Tolita, représentant des pirogues et des pagaveurs.

D'après les cartes marines et les chartes de pilotage modernes, les vents et les courants marins du littoral Nord-équatorial sont favorables à la navigation vers le Nord et défavorables à la navigation vers le Sud : de septembre à novembre, un fort courant (de 15 à 20 nœuds) longe la côte vers le Nord. Il est interdit de faire route vers le Sud, d'autant plus que les vents dominants, même s'ils sont faibles, soufflent aussi vers le Nord quasiment toute l'année. Ce courant faiblit, mais persiste de mars à août. De décembre à février, le courant dominant, orienté Est/Ouest, entraîne vers le large. L'orientation des bancs de sable et l'érosion du littoral façonné par la houle reflètent ces conditions. Longer les côtes vers le Nord entre La Tolita et Tumaco est donc facile, puisqu'on est porté par les courants et par les vents. En revanche, pour voyager en sens inverse, on doit ramer assez vite pour compenser la force adverse des courants et les vents, d'autant plus que la cargaison est lourde et volumineuse, et offre une prise aux vents contraires. D'ailleurs, les navigateurs espagnols évitaient de faire route vers le Sud, tant ils craignaient de ne pas avancer. Ils avaient même inventé l'expression « s'engorgonner » pour signifier que les navires croisant près de l'île de La Gorgona se retrouvaient dans une zone de calme plat et de contre-courants dont ils avaient le plus grand mal à s'échapper. Jusqu'à une époque récente, on pensait que les groupes de la culture Tumaco-La Tolita étaient originaires d'Amérique centrale, mais les fouilles ont montré que leurs ancêtres venaient des régions équatoriales, soit du Sud. Les observations précédentes sur les courants et les vents qui montrent que le voyage par mer était plus facile vers le Nord que vers le Sud, sont en accord avec cette origine méridionale des premiers Indiens de la culture Tumaco-La Tolita.

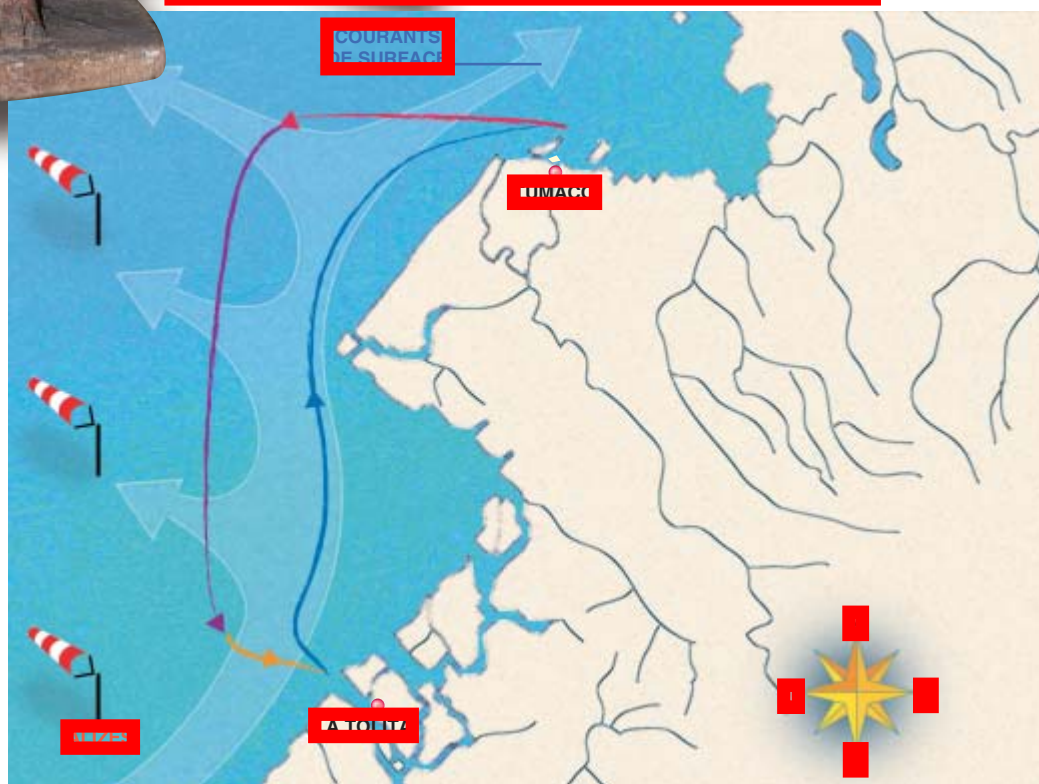
Comment réussissaient-ils à naviguer vers le Sud, notamment pour acheminer leurs morts jusqu'à La

Tolita? Nous supposons qu'ils suivaient une route indirecte, en s'éloignant d'abord de la côte, en direction de l'Ouest, puis qu'ils mettaient le cap au Sud jusqu'à la hauteur du port d'arrivée, et, enfin, qu'ils revenaient vers l'Est (voir la figure 4). Dans cette région de fort marnage, où les courants de flux et de reflux des marées sont puissants près des côtes, ils devaient partir avec la marée descendante et revenir vers la côte avec la marée montante. Pour aller de Tumaco à La Tolita, ils devaient quitter la rade de Tumaco au moment où la marée se retire, et gagner le large avec le flux descendant, puis faire route au Sud pendant l'étape de basse mer et entrer dans l'embouchure du Santiago avec le flux montant. En utilisant ainsi les marées, les pagaveurs devaient

conserver leurs forces pour le segment intermédiaire du trajet, à contre-courant et contre le vent. Cette navigation de cabotage nécessite une longue pratique de la part de marins expérimentés.

Les courants côtiers et la houle expliquent pourquoi les groupes de la culture Tumaco-La Tolita se sont établis sur des îles dans les embouchures des fleuves ou dans des baies abritées plutôt que sur les plages, trop exposées au vent et à la houle (et rares de surcroît). En particulier, le port de transit de l'île d'El Morro, à l'entrée de la rade de Tumaco, face au continent, est bien protégé. De ce port, on peut partir directement vers La Tolita dès que la mer descend.

4. POUR NAVIGUER ENTRE TUMACO ET LA TOLITA, les populations indigènes devaient suivre une route indirecte, car les vents et les courants contraires rendent difficile le voyage par mer vers le Sud. Partant de Tumaco, il profitaient de la marée descendante, s'éloignaient d'abord de la côte (suivant le trajet en rouge), puis pagayaient au large vers le Sud (suivant le trajet en violet), et revenaient ensuite vers la côte avec la marée montante (suivant le trajet en jaune). Au contraire, pour rejoindre Tumaco, ils profitaient du courant et des vents (en bleu) favorables vers le Nord. Ils utilisaient des pirogues taillées dans de gros troncs d'arbre (voir la page de gauche en bas) et mâchaient des chiques de coca (la joue gauche du personnage ci-contre est renflée par une chique) pour lutter contre la fatigue.



Examinons les figurines en céramique de marins payeurs attribuées à la culture Tumaco-La Tolita. L'une des joues présente une boule indiquant que ces payeurs chiquaient (*voir la figure 4*). Les populations traditionnelles préhispaniques mâchaient souvent des feuilles de coca pendant l'effort, et cette pratique s'est perpétuée dans les populations actuelles des Andes. On sait par ailleurs que la coca était utilisée en Équateur dès le troisième millénaire avant notre ère, par les groupes de la culture Valdivia : on a retrouvé des dents présentant des dépôts causés par la chaux qui entrerait dans la préparation de la technique de coca. Par conséquent, les marins consommaient probablement ce stimulant pour pouvoir payer sans arrêt lors du segment de navigation à contre-courant et contre le vent, jusqu'au

moment où le flux montant les propulsait vers la côte.

Si les pirogues naviguaient bien entre Tumaco et La Tolita, quelles marchandises transportaient-elles ? La découverte près de La Tolita de sites agricoles exclut les denrées alimentaires. Les marchandises étaient probablement des biens de grande valeur à la fonction de nécropole de l'île, où un ensemble de rites funéraires étaient pratiqués, étaye cette hypothèse.

La Tolita : nécropole des élites

Dans l'iconographie des objets en or et en céramique de la culture Tumaco-La Tolita, on n'identifie pas de divinités. Les croyances de ces populations précolombiennes s'apparentaient probable-

ment davantage au chamanisme qu'à une religion formée d'un ensemble de croyances structurées avec ses divinités, ses cultes et ses prêtres. Le culte des morts n'en était pas moins d'une grande importance, comme le montrent les sépultures souvent fort riches de la nécropole de La Tolita. Divers objets en or étaient déposés dans les sépultures avec les corps des défunts les plus importants, soit comme offrandes funéraires, soit parce qu'ils étaient des objets personnels. Les autres sépultures de la région ne contenaient pas d'offrandes importantes : La Tolita semble avoir été la seule nécropole où l'on rassemblait les défunts les plus importants. Une grande partie des activités qui s'y déroulaient devait être centrée sur les rites funéraires. L'importance de La Tolita tenait sans doute beaucoup à la présence de prêtres ou de chamanes spécialisés, dont la fonction était d'honorer les défunts par des rites appropriés. Pendant toute la durée de l'hégémonie de La Tolita, sa fonction de nécropole et ses rituels funéraires furent probablement les éléments clés de son influence culturelle, qui touchait sans doute toute la région côtière, depuis Esmeraldas, au Sud, jusqu'à Guapi, voire Buenaventura, au Nord.

Les sépultures se concentrent autour de monticules artificiels (tolas) qui ont donné leur nom à l'île, mais qui ne sont pas des tertres funéraires au sens propre. Ils auraient servi de plate-forme surélevée pour accueillir des édifices ou des espaces cérémoniels à ciel ouvert (qui n'ont pas laissé de trace), dédiés à des pratiques funéraires magiques ou religieuses. Toutefois, certains ont été réutilisés pour creuser des sépultures. La tradition semblait exiger de rassembler les défunts importants dans un endroit particulier, constituant leur terre d'origine culturelle. On leur témoignait ainsi le respect *post mortem* qui leur était dû, étant donné leur statut passé. Les rites funéraires étaient probablement accomplis par un clergé dévoué à cette tâche, qui résidait sur place), et en les isolant sur une île, on protégeait les vivants des pouvoirs que ces défunts auraient pu manifester après leur mort. Pour beaucoup de sociétés andiennes, une forme de présence et de pouvoir invisible du défunt subsistait auprès des vivants. Notamment, le défunt est susceptible de réapparaître sous la forme d'un animal, tel un jaguar ou un autre animal de la forêt. Même s'il était de son vivant



5. LES PERSONNAGES ALLONGES sur le dos et immobilisés par des bandelettes, tel celui-ci, sont fréquents dans l'iconographie des objets en céramique de la culture de Tumaco-La Tolita. Ces figurines semblent représenter les cadavres des défunts, préparés pour leur transport vers la nécropole de La Tolita.

un être bienveillant, il peut devenir maléfique et menacer les vivants si les rituels funéraires n'ont pas neutralisé les pouvoirs qu'il avait avant sa mort. Or, la forêt équatoriale est perçue par ses occupants comme une zone de vigilance obligée où le danger est omniprésent. Avant d'être l'habitat de l'homme, elle est celui d'animaux, souvent féroces, hostiles et capables d'agresser l'homme. Le danger vient aussi de puissances encore plus surnaturelles, qui agissent à travers les félins, les rapaces, les reptiles ou les chauves-souris, et qui pouvaient être perçues comme la manifestation d'un retour des morts parmi les vivants.

En créant une nécropole isolée, dans laquelle les morts resteraient « à leur place », honorés et satisfaits par des rituels funéraires, la société s'assurait une forme de sécurité. Les rites funéraires étaient une sorte de pacte établi avec les puissances surnaturelles des « autres mondes » qui côtoient le monde terrestre ordinaire des vivants.

Outre les innombrables sépultures de La Tolita, les figurines semblent révéler quelques-unes de ces pratiques funéraires. Un grand nombre de figurines anthropomorphes moulées en céramique représentent un personnage allongé sur le dos et immobilisé par des bandelettes ou des sangles (voir la figure 5). Dans le passé, on les a interprétés comme des individus en train d'accomplir des rites d'initiation ou des patients attendant une « opération chirurgicale ». Cependant, l'attitude de ces personnages étendus évoque plutôt des êtres sans vie : nous pensons que ces figurines représentent les cadavres des défunts, préparés pour leur transport vers La Tolita selon la coutume du retour à la terre d'origine culturelle. D'après le caractère unique de la nécropole de La Tolita, on pense que la société organisait pour ses plus importants défunts un « voyage de retour », probablement par voie maritime, au moyen de pirogues creusées dans de gros troncs d'arbre, dont on possède diverses représentations en céramique.

L'occupation de l'île de La Tolita cesse entre 300 et 350. Selon toute probabilité, sa fonction de nécropole ne perdure pas après l'abandon de l'île par ses habitants, mais les populations des alentours en gardent le souvenir. Peu après l'arrivée des premiers Espagnols dans la région, le chroniqueur

Cabello de Balboa rapporte qu'« entre la baie de San Mateo (l'embouchure du rio Esmeraldas) et Ancon de Sardinas, il y a un petit fleuve que les indigènes tiennent pour sacré et lieu de vénération ; ils y viennent pour faire leurs prières et apportent de petits paniers de poudre d'or qu'ils y répandent. Là, séjourne une garnison de 100 Indiens ou plus, qui sont entretenus par les gens de cette province... »

Malgré les siècles qui séparent l'apogée de la culture Tumaco-La Tolita de ce témoignage, l'île reste un lieu sacré, respecté et entretenu par les groupes indigènes contemporains de la conquête, qui conservaient ainsi un lien étroit avec leurs origines.

Toutefois, cette vénération prend fin entre l'arrivée des conquistadores et la visite du moine franciscain Juan de Santa Gertrudis, au XVIII^e siècle, qui nota que l'île de La Tolita était un lieu d'enterrements préhispanique. Il remarque que les monticules artificiels sont érodés par les eaux des grandes marées (ce qui devait déjà être le cas auparavant, mais à une moindre échelle). Il observe aussi que les Indiens de la région y viennent pour récolter, sur les berges du fleuve, divers objets en or ou en céramique que les eaux font apparaître en détruisant les sépultures.

Aujourd'hui, l'érosion et les pillards de sépultures ravagent encore le site de La Tolita. Dans les autres sites archéologiques, la situation est analogue : les pillards détruisent les gisements en quête de figurines, d'ornements ou d'insignes de pouvoir en or, qu'ils revendent, souvent à vil prix, aux intermédiaires alimentant les galeries d'art précolombien en Équateur, en Colombie et à l'étranger. Dans ces pays où la fouille clandestine est une tradition bien ancrée depuis les pillages de la conquête espagnole, la protection du patrimoine culturel et des vestiges précolombiens demeure trop souvent lettre morte, en dépit des efforts des entités locales pour faire respecter le passé et éviter



5. CES BOUCLES D'OREILLES font partie des innombrables objets et figurines en or que le site de La Tolita a livrés. Les défunts les plus importants y étaient enterrés avec des bijoux, des objets personnels et des offrandes faits du précieux métal que les Indiens récoltaient dans les fleuves de la plaine alluviale.

la destruction des sites archéologiques. Par ailleurs, beaucoup de sites, près des mangroves et des cours d'eau, sont menacés de destruction ou ont déjà été détruits par la construction de grands bassins d'élevage d'aquaculture, de plus en plus nombreux dans la région.

Toutefois, le passé précolombien de la région est devenu un peu moins mystérieux depuis que des fouilles archéologiques y ont été entreprises. La longue prospérité de la culture de Tumaco-La Tolita contraste avec les tentatives modernes de développement par la plupart des régions boisées équatoriales, qui, malgré les progrès techniques, n'aboutissent souvent qu'à la destruction du milieu naturel trop fragile pour résister aux agressions de puissants engins mécaniques. Cette culture avait su conserver l'environnement et disposait de ressources bien supérieures à celles des populations actuelles, qui comptent parmi les plus déshéritées du littoral Pacifique équatorial.

J.-F. BOUCHARD, archéologue au CNRS et directeur des missions françaises en Colombie et en Équateur de 1976 à 2000.

A. ALCINA-FRANCH, *La arqueología de Esmeraldas (Ecuador). Introducción general. Memorias de la Misión Arqueológica Española en el Ecuador*, vol. 1, Ministerio de Asuntos Exteriores, Madrid, 1979.

J.-F. Bouchard, *Recherches archéologiques*

dans la région de Tumaco (Colombie), in *Mémoire*, n° 34, Institut français d'études andines, Éditions Recherche sur les Civilisations, A.D.P.E., Paris, 1984.

J.-F. BOUCHARD, *Occupations préhispaniques à El Morro-Tumaco (Colombie)*, Colloque d'archéologie des Andes équatoriales, département d'ethnologie, Univ. de Neuchâtel, Bulletin de la Société suisse des américanistes, n° 63, pp. 111-116, février 1999.

L'enfer des paris

JEAN-PAUL DELAHAYE

Certaines formulations sont parfois si simples qu'on croit pouvoir maîtriser les probabilités. À tort : il faut nous méfier de l'intuition.

Les probabilités sont un domaine où l'intuition nous joue des tours. Nous allons présenter trois jeux étranges où notre intuition a du mal à produire un jugement correct, ce qui conduit à des situations que nous jugeons paradoxales. La première est d'autant plus étonnante qu'elle est simple, la deuxième est déjà plus subtile, la troisième laisse carrément perplexe.

S'OPPOSER AU HASARD DES NAISSANCES?

Dans un pays lointain, les femmes ont des enfants qui sont de sexe masculin dans 50 pour cent des cas exactement et, bien sûr, de sexe féminin dans 50 pour cent des cas. Le gouvernement décide que seuls les couples ayant eu au moins une fille toucheront leur retraite. En réaction à cette mesure, chaque couple adopte alors la stratégie suivante :

- si leur premier enfant est une fille, satisfait il n'en a pas d'autres ;
- si le premier enfant est un garçon, le couple a un second enfant qui sera le dernier si c'est une fille ;
- et ainsi de suite, chaque couple ayant des enfants jusqu'à ce qu'il ait une fille laquelle est alors leur dernier enfant.

Cette stratégie favorise les filles puisque les familles s'arrêtent dès qu'une fille naît. Sans faire de calcul et en réfléchissant une minute au plus, donnez votre avis : la proportion de filles à la naissance dans ce pays est-elle de $1/3$, $1/2$, $2/3$, $3/4$ ou $4/5$? (l'une des réponses proposées est juste)

Réponse

La seconde réponse de la liste, $1/2$, est la bonne : les stratégies des familles n'ont aucune influence sur la proportion de filles (il en serait de même si la probabilité de naissance des filles était différente de 50 pour cent). Aussi surprenant que cela paraisse, les stratégies familiales appliquées par les couples n'ont aucun effet sur le rapport garçon/fille. On peut s'en rendre compte sans faire le moindre calcul : quelle que soit la position d'une fille dans une famille, la probabilité que le prochain

enfant soit de sexe féminin est $1/2$: le passé n'influe pas sur la naissance à venir (c'est du moins l'hypothèse raisonnable qu'on adopte toujours). Tout enfant à naître ayant une probabilité de 50 pour cent d'être une fille, il naît donc en moyenne une fille pour deux naissances. Plus généralement, quelles que soient les règles adoptées par les familles pour cesser d'avoir des enfants en fonction des précédentes naissances dans la famille, les proportions de filles et de garçons restent inchangées : toutes les stratégies se valent et aucune n'a le moindre effet perturbateur.

Auriez-vous des doutes concernant le raisonnement ci-dessus ? Faites alors le calcul suivant (pour simplifier nous supposons que les familles n'ont jamais plus de quatre enfants, mais vous pouvez reprendre le calcul avec 5, 6 ou n enfants, ou même sans limitation du nombre d'enfants).

Probabilité qu'une famille possède un seul enfant : $1/2$. La famille est alors du type : [fille].

Probabilité qu'une famille possède deux enfants : $1/4$. La famille est du type : [garçon, fille].

Probabilité qu'une famille possède trois enfants : $1/8$. La famille est du type : [garçon, garçon, fille].

Probabilité qu'une famille possède quatre enfants : $1/8$. La famille est une fois sur deux du type : [garçon, garçon, garçon, fille] et une fois sur deux du type : [garçon, garçon, garçon, garçon].

Donc, sur 16 familles, il y a en moyenne huit familles du type [fille], quatre du type [garçon, fille], deux du type [garçon, garçon, fille], une du type [garçon, garçon, garçon, fille] et une du type [garçon, garçon, garçon, garçon]. Cela fait au total $8 + 4 + 2 + 1 = 15$ filles et $4 + 4 + 3 + 4 = 15$ garçons.

L'effort fait par chaque famille pour avoir une fille ne change pas la proportion de garçons et de filles, mais conduit cependant à une situation où la plupart des familles sont satisfaites, car elles ont au moins une fille (dans notre exemple 15 familles sur 16 ont une fille).

Signalons hélas, qu'aujourd'hui en Inde et en Chine, les règles traditionnelles

sur les dots et d'autres raisons socio-culturelles font que les familles souhaitent avoir en priorité des garçons. Il en résulte que la proportion de filles à la naissance est nettement inférieure à celle des garçons. Cela n'est pas la conséquence de stratégies du type de celle envisagée plus haut, mais cela est dû à des avortements sélectifs organisés par les familles qui, grâce aux échographies, savent au bout de quelques semaines de grossesse le sexe de l'enfant à naître. Ce type de comportement est combattu par les autorités, car il conduit à un déséquilibre entre hommes et femmes susceptible à terme de créer des problèmes sociaux. Dans plusieurs régions d'Inde on compte déjà plus de 110 garçons pour 100 filles et on prévoit que le problème va s'aggraver.

LE ROUGE EST MIS...

Dans le problème précédent, les «tirages» étaient indépendants et c'était l'explication simple de l'impossibilité d'adopter des stratégies familiales modifiant la proportion globale d'hommes et de femmes. Nous allons aborder un problème de stratégie optimale plus intéressant (dû au mathématicien prestidigitateur P. Diaconis) où, au contraire, les événements dépendent fortement du passé.

On vous propose le jeu suivant dénommé *La prochaine est rouge* :
 - on bat les cartes d'un paquet de 32 cartes ;
 - le meneur de jeu retourne les cartes une à une ;
 - à un moment librement choisi par vous, vous arrêtez le meneur de jeu et vous annoncez : «La prochaine carte sera rouge» :

- le meneur retourne la carte, si vous avez raison vous avez gagné, sinon vous avez perdu.

Vous pouvez utiliser la méthode que vous voulez et, par exemple, attendre que presque toutes les cartes soient passées. Vous pouvez même attendre qu'il ne reste plus qu'une seule carte. Notez bien cependant que vous ne pouvez pas choisir entre rouge et noire (sinon, en attendant qu'il ne reste qu'une carte, vous gagneriez à chaque fois) : vous êtes obligé de parier sur rouge.

La stratégie évidente consistant à attendre la dernière carte vous garantit de gagner une fois sur deux au moins. La stratégie consistant à parier dès la première carte donne aussi une chance sur deux de gagner. Il ne faut donc ni se précipiter ni trop attendre pour avoir de bonnes chances.

La question est :

- on vous propose de jouer à ce jeu à contre 2 : si vous gagnez, le meneur de jeu vous donne un franc, si vous perdez, vous donnez deux francs au meneur de jeu. Devez-vous accepter, et si oui quelle méthode de jeu devez-vous adopter pour maximiser votre gain ? Une façon équivalente de poser le problème est : existe-t-il une stratégie de prédiction qui vous garantira de gagner deux fois plus souvent que vous ne perdrez ?

Réponse

La réponse est négative, mais elle est plus étonnante encore, car, quelle que soit la stratégie que vous utiliserez, vous ne pourrez jamais avoir plus d'une chance sur deux de gagner et vous ne pourrez d'ailleurs jamais non plus avoir moins d'une chance sur deux de gagner. La stratégie courte consistant à parier dès

la première carte est aussi bonne que celle consistant à attendre la dernière et cette stratégie est équivalente à toute autre stratégie aussi compliquée ou subtile que vous pourriez imaginer.

Comme dans le problème des naissances et malgré le contexte différent d'un paquet de cartes fini qui s'épuise en vous procurant des informations sur ce qu'il reste à venir, vous ne pouvez pas changer le 1/2 : vous ne pouvez ni faire mieux ni moins bien !

Un raisonnement qui montre l'équivalence de toutes les stratégies a été récemment proposé par Michel Emery de l'Université de Strasbourg. Il est présenté sur la figure 2.

MÉFIEZ-VOUS DES PARIS GROUPÉS

Voici pour terminer un extraordinaire paradoxe concernant les paris groupés. Pour bien en saisir la portée, examinons ce qu'est un pari groupé.

On joue avec un dé à six faces non truqué. Considérons les deux paris suivants, qui vous sont tous les deux favorables :

- Pari 1 : si c'est 6, tu gagnes 2 F, si c'est 1 tu perds 1 F, sinon rien.
- Pari 2 : si c'est 5, tu gagnes 2 F, si c'est 2 tu perds 1 F, sinon rien.

En acceptant les deux paris simultanément, vous gagnerez en moyenne deux fois plus que vous ne perdrez ; ce sera donc toujours un pari favorable pour vous.

D'une manière générale, il est clair que si l'on vous propose simultanément plusieurs paris qui, chacun pris individuellement vous sont favorables, cela constitue un nouveau pari (que nous nommerons le pari groupé). Lequel vous est toujours

PEUT-ON MODIFIER LE RAPPORT GARÇONS/FILLES ?

Même si tous les couples adoptaient la règle de ne plus avoir d'enfants dès qu'une fille naît, ce qui semble favoriser les filles, cela n'aurait aucun effet sur le rapport garçons/fille. Montrons-le pour le cas (idéalisé) où chaque couple a des enfants indéfiniment (sans limitation du nombre) jusqu'à ce qu'une fille naisse.

Probabilité qu'une famille possède un seul enfant : 1/2.

La famille est du type : [fille]

Probabilité qu'une famille possède 2 enfants : 1/4. La

famille est du type : [garçon, fille]

Probabilité qu'une famille possède 3 enfants : 1/8. La

famille est du type : [garçon, garçon, fille]

..

Probabilité qu'une famille possède n enfants : $1/2^n$. La

famille est du type : [garçon, garçon, ..., garçon, fille]

Calculons la proportion de garçons et de filles en ne comptant que les familles de n enfants ou moins (on fera ensuite tendre n vers l'infini).

Sur $2^n - 1$ familles, il y en aura en moyenne 2^{n-1} du type [fille], 2^{n-2} du type [garçon, fille], 2^{n-3} du type [garçon, garçon, fille]... une du type [garçon, garçon, ..., garçon, fille], etc.

Il y a donc, au total, $2^{n-1} + 2^{n-2} + \dots + 1 = 2^n - 1$ filles.

Pour compter les garçons, le travail est un peu plus compliqué.

Nombre d'aînés : $2^{n-1} - 1$

Nombre de seconds : $2^{n-2} - 1$

Nombre de troisièmes : $2^{n-3} - 1$

..

..

Nombre de $(n - 1)$ -ièmes : $2 - 1$

Nombre de n -ièmes : $1 - 1$

Au total, le nombre de garçons est donc :

$2^{n-1} + 2^{n-2} + 2^{n-3} + \dots + 2 + 1 - n = 2^{n-1} - 1 - n$.

Le rapport filles/garçons a bien pour limite 1 à l'infini.

NOEL

POUR VOUS GAGNER UNE FOIS SUR DEUX, IL NE VOUS RESTE PLUS NI MOINS

À FAIRE : LA PROCHAINE EST ROUGE

Imaginons un grand tableau composé de :

$32! = 26313083693369353016721801216000000$ lignes,

chacune contenant un classement possible d'un jeu de 32 cartes.

Une stratégie consiste à fixer une règle du type «lorsque la dixième carte noire est passée, je parie que la suivante est rouge». Une stratégie ne dépend que du passé et détermine, en fonction de celui-ci, si j'arrête le meneur de jeu en lui disant «la prochaine est rouge».

Fixons une méthode de jeu. Dans chaque ligne du tableau, nous plaçons une étoile marquant l'endroit du pari, ce qui nous donne donc un tableau de $32!$ lignes, chacune contenant une étoile. Voici un exemple de ligne : l'étoile * est placée juste après la dixième carte noire et fait gagner le joueur, car une carte rouge (le R♦) suit) :

V♣, D♥, V♥, 10♠, R♠, As♠, 8♦, 8♥, 7♣, 10♦, As♥, 8♣, 9♦, V♠, D♠, D♠, As♦, 9♥, R♠, *, R♦, 7♠, 8♠, 9♣, 9♠, 10♣, As♠, V♦, D♦, R♥, 7♥, 10♥, 7♦. Nous devons évaluer le nombre de lignes où l'étoile * est suivie

d'une carte rouge (partie gagnante pour le parieur) et le nombre de lignes où elle est suivie d'une carte noire (partie perdante pour le parieur).

Nous allons montrer que ces deux nombres sont égaux. Pour cela, nous opérons la transformation suivante de notre tableau : dans chaque ligne, nous permutons la carte située après l'étoile et la dernière carte.

Ainsi la ligne xx... xx*Ayy... yyB devient xx... xx*Byy... yyA

Cette transformation revient à changer l'ordre de deux lignes du tableau.

En effet, le tableau contient toutes les permutations possibles des 32 cartes, et échanger deux cartes revient à échanger deux lignes. Dans le cas où l'étoile est juste avant la dernière carte (le pari a été fait au dernier moment et la méthode permet ce choix), la ligne xx... xx*A reste identique. Au total, le tableau transformé est donc constitué des $32!$ lignes du tableau initial ; seul l'ordre des lignes a été changé.

Dans le tableau initial, il y avait exactement le même nombre de lignes se terminant par une carte noire que de lignes se terminant par une carte rouge (les rouges et les noirs ont des rôles symétriques), donc dans le tableau transformé aussi. Et donc, il a exactement le même nombre de lignes se terminant pour une carte A de couleur rouge que de lignes se terminant par une carte A de couleur noire, ce qui signifie que la stratégie fait gagner le parieur exactement une fois sur deux. Notre raisonnement ne dépend pas de la stratégie considérée et donc, quelle que soit la stratégie utilisée pour le jeu «La prochaine est rouge», elle donne une chance sur deux de gagner, ni plus ni moins.

Si vous ne croyez pas en ce raisonnement, ou qu'il vous semble faux, alors vous restez persuadé qu'au jeu «La prochaine est rouge», vous pouvez gagner plus d'une fois sur deux. Dans ce cas, prenez contact avec moi, je suis prêt à jouer le meneur de jeux à 1 contre 2 (si vous avez raison, vous gagnez un, sinon vous me donnez deux). J'accepte aussi de jouer avec 10 cartes au lieu de 32.

favorable. En bref, on a ce qu'on pourrait appeler la loi des paris groupés :

- un pari groupé de paris favorables est un pari favorable.

On peut même remarquer que le gain moyen (ou espérance mathématique) d'un pari groupé est la somme des gains moyens de chaque pari. Ici, le gain moyen pour chaque pari est de $1/6$ (une fois sur six, on gagne 2 F, une fois sur 6, on perd 1 F. Donc sur six coups on gagne en moyenne 1 F, ce qui donne un gain moyen de $1/6$ par pari). Le gain moyen pour le pari groupé est de $2/6$ soit un tiers de franc.

Très étrangement, le mathématicien Leonid Levin, célèbre pour de nombreux travaux réalisés en théorie de la complexité, a trouvé une situation qui contredit la loi des paris groupés.

Décrire la situation va être un peu compliqué, mais le jeu en vaut la chandelle, car on a affaire à un paradoxe qui, contrairement aux deux précédents, ne semble pas avoir été résolu de façon satisfaisante. Le paradoxe est fondé sur les nombres.

Tout le monde connaît les nombres décimaux, qui sont ceux qui s'écrivent de manière finie en base 10, comme

0,132 ; -3,3211 ; 0,00012 ; -219,975,54321 ; etc. Nous noterons D l'ensemble de tous les nombres décimaux positifs ou négatifs).

L'ensemble des nombres réels (cette fois, les développements décimaux peuvent être finis ou infinis) peut être séparé en classes décimales : on constitue des paquets de nombres réels en mettant ensemble tous ceux dont la différence est un nombre décimal.

Par exemple, on a une classe

$\pi : \pi + 0,132 : \pi - 3,3211$; etc.

Une autre classe on mettra

$\sqrt{2} : \sqrt{2} + 0,132 : \sqrt{2} - 3,3211$; etc.

Il y a une infinité non dénombrable de classes décimales (l'argument mathématique est : si ce n'était pas le cas, l'ensemble des nombres réels serait dénombrable, ce qui est faux).

Tout nombre réel appartient à une classe décimale et à une seule. Par commodité, on choisit un chef dans chaque classe décimale, ce chef (ou représentant de la classe) permettant de donner un nom à la classe. La classe dont le chef est π sera notée Classe(π). Par définition, Classe(π) est la classe de tous les nombres réels de la forme $\pi + d$, avec d qui est un nombre décimal.

Puisqu'il y a une infinité de classes, il y a une infinité de chefs de classe. On note C l'ensemble de tous les chefs de classe ; D est l'ensemble des nombres décimaux. Les deux ensembles C et D ont la propriété intéressante suivante, qui résulte de leur définition : tout nombre réel s'écrit de manière unique $c + d$ avec c dans C et d dans D (c est le chef de la classe à laquelle appartient x , et d est le nombre décimal $x - c$).

Considérons maintenant le pari suivant :

- on choisit un nombre réel x au hasard entre 0 et 1. Pour cela, on choisit un point quelconque sur un segment de longueur 1 et on mesure sa distance à une extrémité choisie comme origine, ou, ce qui revient au même, on lance un dé à 10 faces indéfiniment et on note tous les chiffres obtenus qui déterminent un nombre réel comme 0,271828182845...)
- le nombre tiré x s'écrit sous la forme $c + d$, avec c , un chef, et d , un décimal ;
- si le nombre choisi est de la forme $c + 0,7$, vous gagnez 2 F ;
- si le nombre choisi est de la forme $c + 0,9$, vous perdez 1 F ;
- sinon, personne n'a gagné, personne n'a perdu.

Un tel pari vous est favorable, car vous gagnerez aussi souvent que vous perdrez, mais lorsque vous gagnerez, vous gagnerez deux fois plus que lorsque vous perdrez (comme tout à l'heure avec le dé à six faces). Plus généralement soit c et d' deux nombres décimaux quelconques

nous considérerons le pari suivant que nous noterons $P(d; d')$:

- nous choisissons un nombre réel x au hasard entre 0 et 1. Nous déterminons sa classe c ;
- si le nombre choisi est de la forme $c + d$ vous gagnez 2 F ;
- si le nombre choisi est de la forme $c + d$ vous perdez 1 F ;
- sinon, personne n'a gagné, personne n'a perdu.

Si l'on vous propose le pari $P(d; d')$, vous devez l'accepter, car le nombre x choisi a autant de chances d'être de la forme $c + d$ que de la forme $c + d'$. Comme précédemment, le pari vous est favorable, car vous gagnez plus que vous ne risquez.

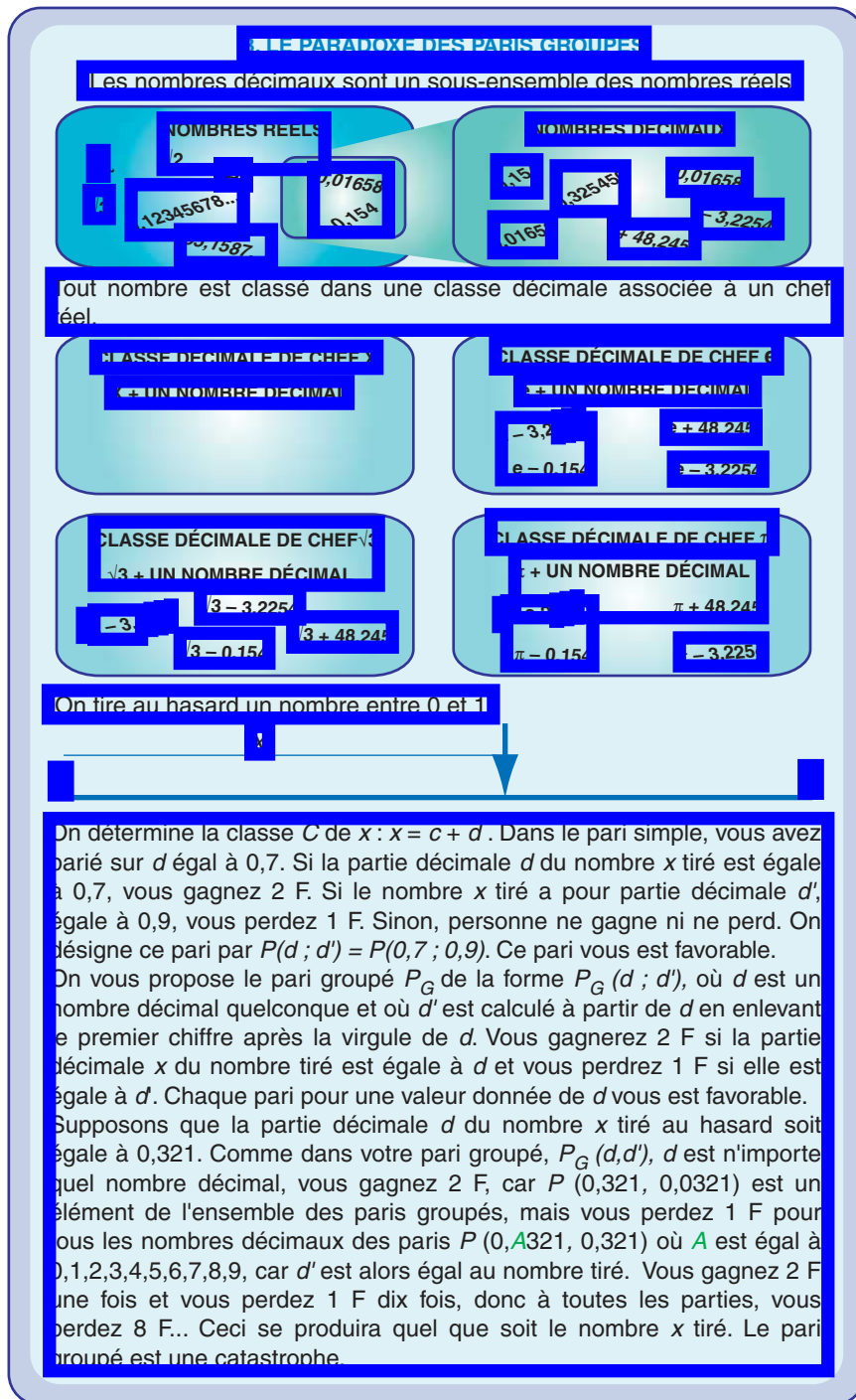
On peut détailler l'argument indiquant que x a autant de chances d'être de la forme $c + d$ que de la forme $c + d'$. Si on représente sur un cercle de périmètre 1, muni d'une origine l'ensemble des cas qui vous sont favorables (tous les nombres réels entre 0 et 1 de la forme $c + d$ avec c un chef), et qu'on fait tourner cet ensemble de $d' - d$, on obtient exactement l'ensemble des cas qui vous sont défavorables (les nombres de la forme $c + d'$). L'ensemble des cas défavorables s'obtenant par rotation de l'ensemble des cas favorables, en choisissant un point au hasard sur le cercle, on a bien autant de chances de tomber sur un cas favorable que sur un cas défavorable. Comme les cas favorables rapportent le double de ce que coûtent les cas défavorables, il faut accepter le pari.

On vous propose maintenant un pari groupé, composé de paris du type précédent. Précisément, on vous propose le pari $P(0,235 ; 0,35)$ en même temps que le pari $P(0,78432 ; 0,8432)$ et tous les paris de la forme $P(d,d')$, où d' est obtenu en supprimant le premier chiffre après la virgule de d . Ce pari groupé est des paris favorables. Vous brûlez d'en faire l'accepter.

Pourtant, si vous acceptez ce par
groupé, vous allez vous en mordre les
doigts car, très étrangement, il vous es
ortement défavorable.

En effet : supposons que le nombre réel tiré x soit de la forme $\pi + 0,321$. Faisons le bilan des paris gagnés et perdus : vous avez gagné le pari $P(0,321 ; 0,21)$ ce qui vous rapporte 2 F ; vous avez perdu le pari $P(0,0321 ; 0,321)$, ce qui vous coûte 1 F ; vous avez perdu le pari $P(0,1321 ; 0,321)$, ce qui vous coûte 1 F ; etc. pour tous les paris $P(0,2321 ; 0,321)$, $P(0,3321 ; 0,321)$, ..., $P(0,9321 ; 0,321)$.

Au total vous avez perdu dix paris (donc 10 F) et vous avez gagné 2 F. Bilan : le pari groupé vous coûte 8 F à chaque fois. Bien sûr, quel que soit le nombre x



choisi la situation est semblable et vous perdez donc 8F à chaque fois. (Dans le cas où le nombre choisi au hasard est 0 la situation est particulière : vous ne gagnez même pas 2 F et perdez donc 10 F d'un coup.)

Comment expliquer qu'un pari groupé de paris favorables vous soit défavorable ?

Réponses

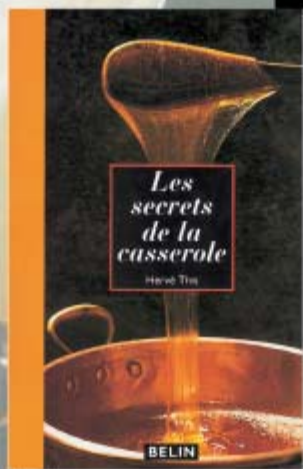
En fait, je ne connais pour l'instant aucune explication totalement satisfaisante au paradoxe de Levin. Je vais donc fournir plusieurs pistes avec à chaque fois une critique et je suggère

aux lecteurs de *Pour la science* de me
donner leur avis.

(a) Première possibilité d'explication
l'axiome du choix

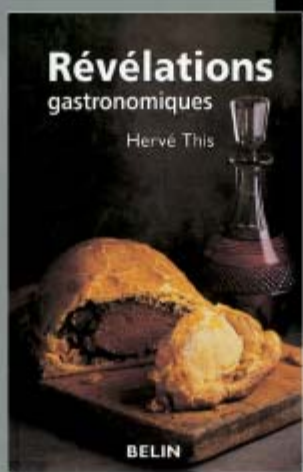
Il faut remarquer que nous utilisons l'axiome du choix. Lorsque nous disons que nous pouvons prendre un chef dans chaque classe, nous sommes en train d'utiliser l'axiome du choix. Cet axiome indique que lorsqu'on a un ensemble d'ensembles non vides (ici les classes décimales) il est possible de prendre un élément de chaque ensemble et de les regrouper pour faire un nouvel ensemble (c'est de cette façon

Une nouvelle rencontre de la science et de l'art.



LES SECRETS DE LA CASSEROLE

Comment préparer plus de 20 litres de mayonnaise à partir d'un seul jaune d'œuf ? En explorant la physique des émulsions. Comment obtenir un bon rôti ? En empruntant à la chimie ses études des sucres et des acides aminés. Comment obtenir de bonnes gelées ? En reprenant les résultats récents de la physico-chimie des gels. Code 001585 - 130 F-19,82 € 224 pages



RÉVÉLATIONS GASTRONOMIQUES

Encore un livre de recettes ! Oui, mais des recettes garanties par des expérimentations rigoureuses, par des analyses physico-chimiques des aliments et des tours de main culinaires. Chaque geste culinaire est ici disséqué, commenté du triple point de vue culinaire, physico-chimique et gastronomique. Code 001756 - 130 F-19,82 € 320 pages

Ces deux ouvrages ont obtenu le Prix de l'Académie nationale de Cuisine 1997.

LA CASSEROLE DES ENFANTS

Laissés seuls pour la soirée, deux enfants doivent cuisiner leur dîner. Mais leur sens aigu de l'observation et de l'expérimentation leur fait découvrir les bases de la chimie et de la physique. Ce livre pourrait être un manuel, s'il n'évoquait les molécules et leurs transformations de façon si simple et si amusante.

Code 002309 - 80 F - 12,20 € - 128 pages



BELIN

Voir bon de commande page 96

que nous avons constitué l'ensemble C (des chefs de classe). On sait que l'axiome du choix engendre des situations paradoxales, et ce qui se passe ici n'est qu'un exemple nouveau de ces situations paradoxales provoquées par l'axiome du choix.

Critique : l'axiome du choix engendre des situations choquantes pour l'intuition, mais jamais de vraies contradictions. Ici, il est tellement absurde qu'un groupement de paris favorables devienne défavorable qu'on a affaire à un problème de grave incohérence. Notons aussi qu'habituellement, l'axiome du choix est accepté. Doit-on décider de ne plus l'utiliser en probabilité ?

(b) *Deuxième possibilité d'explication de la loi des paris groupés.*

La règle qu'un pari groupé composé de paris favorables est favorable n'est peut-être pas vraie en général. Notre intuition nous souffle cette loi, mais elle se rompt peut-être, comme cela lui arrive.

Critique : on prouve sans mal qu'un pari groupé composé d'un nombre fini de paris favorables est favorable. Précisément, l'espérance de gain d'un pari groupé est la somme des espérances de gain des paris individuels. Dans le cas infini la situation est la même et une somme infinie de gains moyens positifs est bien évidemment positive. La loi des paris groupés est bonne et on ne peut la refuser.

(c) *Troisième possibilité d'explication des probabilités continues.*

Dans le pari considéré, on tire un nombre réel au hasard. Or cela n'est pas possible en pratique et d'ailleurs, concrètement, on ne pourrait pas mettre en œuvre le pari comme c'est suggéré en prenant un point au hasard sur une droite ou en lançant une infinité de fois un dé à dix faces.

Critique : en théorie des probabilités on considère sans arrêt des situations analogues et la théorie de la mesure donne un sens précis à l'idée d'un nombre réel tiré au hasard entre 0 et 1. Si ce type de tirage continu n'a pas de sens, pourquoi tant de livres lui sont-ils consacrés ? Et s'il a un sens habituellement (et est utile concrètement dans de nombreuses applications), pourquoi ici n'en aurait-il pas ?

Je supplie donc les lecteurs de m'indiquer une réponse plus satisfaisante.

A Bet with Leonid Levin, in *The Mathematical Intelligencer*, vol. 23, n° 1, p. 42 2001.

Michel EMERY, *Quelques phénomènes curieux en probabilités et statistiques* in *L'Ouvrier*. 2001. article à paraître.

Entrez dans le quadrille

DENNIS SHASHA

Une chorégraphe sollicite votre aide pour son prochain spectacle

Concentré sur la recherche de la solution de l'énigme du mois dernier, vous ne voyez pas arriver dans votre bureau la chorégraphe de la célèbre compagnie «102 sous 2» qui a fait un triomphe la saison passée. Elle a besoin de vos talents pour résoudre un problème qui entrave la mise au point de son prochain ballet. Après lui avoir juré de garder le secret le plus absolu sur votre (éventuelle) participation, elle vous expose la situation.

«Ma compagnie est composée de 12 hommes (les carrés bleus sur la figure) et de 20 femmes (en rose). Avant le final, les hommes encerclent les femmes selon cette disposition.»

Elle dessine une configuration (voir la figure ci-contre) sur le verso d'une feuille sans se soucier du recto, puis poursuit :

«Je veux qu'après seulement trois mouvements, les femmes encerclent les hommes de cette façon.»

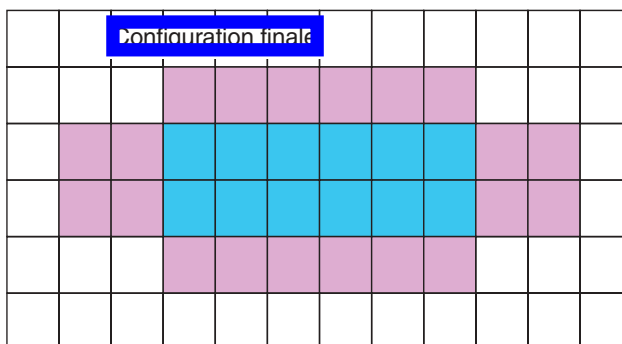
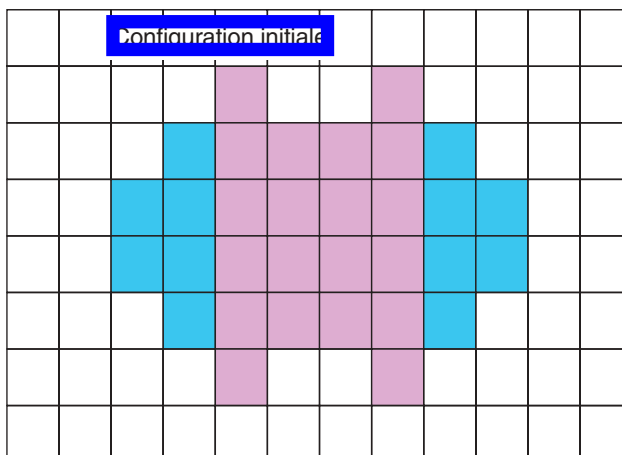
Elle dessine alors la deuxième configuration.

«Pendant chaque mouvement, un danseur peut rester sur place ou se déplacer d'un pas dans l'une des quatre directions : à gauche, à droite, en avant ou en arrière. J'impose deux conditions. D'une part, deux danseurs ne peuvent s'échanger leur place pendant un mouvement, d'autre part, deux danseurs ne peuvent être à la même place à la fin d'un mouvement. J'évite toute collision.»

«Vous avez les configurations initiale et finale, déterminez moi les déplacements de chaque danseur à chaque mouvement!»

Son exposé terminé, elle s'assied, bien décidée à ne parler qu'avec la solution. À vous de l'aider... Pour ce faire, pensez à vous munir d'un papier quadrillé.

La solution de l'énigme sera donnée le mois prochain, et sera bientôt disponible sur le site Internet : www.pourlascience.com



Le problème d'énigmaths du numéro 286

Pour le premier problème, vous pouvez espérer gagner 9 309,091 euros. Définissons $S(x, t, n)$, la somme dont vous disposez à la fin de la série de paris quand vous avez x euros, quand il reste t jets de pièce et que l'oracle a n mensonges à sa disposition (n est égal à 1 ou à 0).

Lorsque vous avez à jouer t fois, que l'oracle ne peut plus mentir et que vous avez x euros, vous pariez tout votre argent à chaque fois : $S(x, t, 0) = x \cdot 2^t$. Dans le cas où il reste un pari et que l'oracle peut encore mentir, vous ne pariez rien : $S(x, 1, 1) = x$.

Calculons $S(x, t+1, 1)$. Si vous mettez en jeu b euros au pari suivant, soit l'oracle ment et vous finirez la partie avec $S(x - b, t, 0)$, soit l'oracle ne ment pas cette fois et vous aurez $S(x + b, t, 0)$. Opti-

miser le gain consiste à évaluer $S(x - b, t, 0)$ et $S(x + b, t, 0)$, car l'oracle vous fera choisir la pire solution. Étudions l'équation $S(x - b, t, 0) = S(x + b, t, 0)$.

$S(x, 1, 1) = x$, $S(x, 1, 0) = 2x$, $S(x, 2, 0) = x^2$, $S(x, 2, 1) = 4x/3$ (vous avez parié $x/3$), $S(x, 3, 1) = 2x$ (vous avez parié $x/2$). Pour $S(x, 4, 1)$, on résout $S(x - b, 3, 0) = S(x + b, 3, 1)$. Soit $8(x - b) = 2(x + b)$. $b = 6x/10$ et $S(x, 4, 1) = 32x/10$. Et ainsi de suite...

Ainsi, $S(100, 2, 1) = 133,33$ euros.

$S(100, 3, 1) = 200$ euros,
 $S(100, 4, 1) = 320$ euros,
 $S(100, 5, 1) = 533,33$ euros,
 $S(100, 6, 1) = 914,28$ euros,
 $S(100, 7, 1) = 1600$ euros,
 $S(100, 8, 1) = 2844,44$ euros,
 $S(100, 9, 1) = 5120$ euros,
 $S(100, 10, 1) = 9309,09$ euros.

$S(100, 11, 1) = 17066,67$ euros. Après 20 jets de pièce, vous auriez plusieurs millions d'euros!

Pour le second problème, quand vous devez parier à l'avance, vous n'aurez au final que 1 600 euros au mieux. Créons de nouveau deux fonctions. $M(n)$ est le montant minimal que vous êtes sûr d'avoir après n jets si le mensonge a déjà eu lieu. $M(n) = 100 +$ les sommes des $n - 1$ paris - la somme du pari le plus élevé.

$V(n)$ est la somme que vous avez après n jets quand le mensonge est encore à venir. $V(n) = 100 +$ tous les paris gagnés.

La meilleure stratégie qui se présente est de parier successivement : 50 ; 50 ; 100 ; 100 ; 200 ; 200 ; 400 ; 400 ; 800 ; 800. Ce qui se traduit par les gains successifs suivants : 50 ; 100 ; 100 ; 200 ; 200 ; 400 ; 400 ; 800 ; 800 ; 1600.

Les moirés de Vasarély

LOIC MANGIN

L'œuvre de Vasarély nous invite à une promenade entre la physique et les mathématiques.

«Un mathématicien est un faiseur de motifs. Ces motifs, comme ceux des peintres ou des poètes, doivent être beaux ; les idées, comme les couleurs ou les mots, doivent s'assembler harmonieusement.»

Gottfried Harms

En 1965, au musée d'Art moderne de New York, sont regroupées pour la première fois les œuvres dites d'art optique, car elles jouent sur les illusions perceptives, les effets de moirés, les scintillements... Parmi les artistes que l'on y admire figurent Francis Morelet, et surtout l'initiateur de cette forme d'art, le peintre français d'origine hongroise Victor Vasarély (1906-1997). Par son succès auprès des critiques et du public, Vasarély est rapidement baptisé «l'inventeur de l'art optique». À cette époque, l'engouement est tel que de nombreuses revues, de *Vogue* à *Scientific American*, font leur couverture des œuvres d'art optique. Les uns pour leur esthétique, les autres pour leur richesse géométrique.

Victor Vasarély a fréquenté le Műhely de Budapest, l'équivalent hongrois du Bauhaus, l'école d'art fondée à Weimar en 1919 par Walter Gropius. Pour ce dernier, le but de toute activité plastique est la construction et l'architecture. Vasarély a ainsi suivi, indirectement, les enseignements des professeurs de l'école disparue en 1933, tels Josef Albers (1888-1976), Paul Klee (1879-1940) ou Wassily Kandinsky (1866-1944), les inventeurs de l'art abstrait.

Installé à Paris en 1930, il devient l'un des maîtres de l'abstraction géométrique et concrétise dans ses œuvres ses recherches théoriques, notamment celles sur la perspective. En 1953, dans une série nommée *Transparence*, il joue avec les effets de moirés (voir la figure 3), où un motif périodique de lignes est couvert du même motif (ou de son négatif) imprimé sur un support transparent. De petits déplacements de ce dernier créent les motifs moirés. Survolons les bases mathématiques de ces figures, ainsi que les services qu'elles rendent aux physiciens.

Le moiré le plus simple apparaît quand deux séries identiques de lignes verticales régulièrement espacées sont superposées avec un léger décalage ou une rotation de quelques degrés. Des larges rayures sombres apparaissent : plus l'alignement est parfait, plus ces rayures sont espacées.

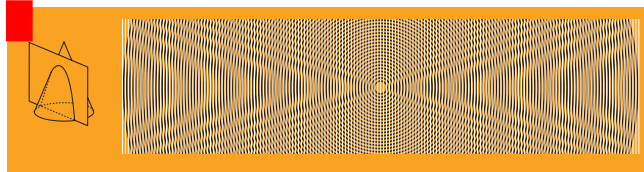
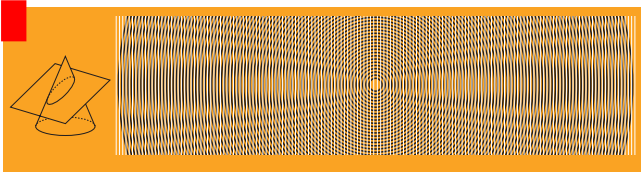
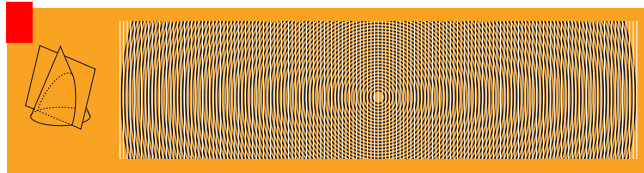
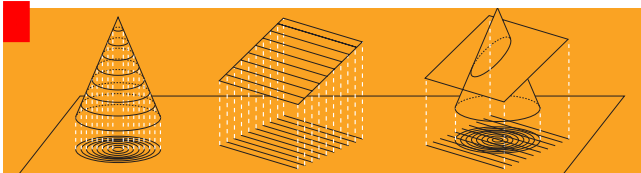
Pour comprendre l'origine des figures de moirés, considérons-les en termes de géométrie projective. Ce domaine des mathématiques est celui

des objets d'un espace à trois dimensions projetés sur un espace à deux dimensions. Par exemple, quand on projette les intersections d'un cône avec une série de plans, parallèles à sa base et équidistants, on obtient sur un plan de projection parallèle à la base du cône une série de cercles concentriques (voir la figure 1).

De la même façon, les intersections d'un plan incliné par une série de plans perpendiculaires au plan de projection se projettent en une série de lignes parallèles. Ces lignes parallèles sont d'autant plus écartées que le plan initial est incliné par rapport au plan de projection.

Ainsi, en géométrie projective, un cône est équivalent à une série de cercles concentriques et un plan, à une grille. Quels motifs moirés apparaissent lors de la superposition des cercles concentriques et de la grille ? Des ellipses, des hyperboles, des paraboles (voir la figure 1), c'est-à-dire les différentes intersections d'un plan et d'un cône, ce que les mathématiciens nomment les sections coniques.

Quand l'espacement de la grille est environ égal à celui des cercles concentriques, les figures de moirés obtenues sont des paraboles (1b). Lorsque l'espacement augmente, c'est-à-dire que le plan initial est de plus en plus hori-



1. Les figures de moirés résultent de la superposition de deux réseaux de cercles ou de droites. Quand on projette les intersections d'un cône par une série de plans parallèles au plan de la base (a, à gauche), on obtient des cercles concentriques. Par ailleurs, les intersections d'un plan incliné par une série de plans perpendiculaires au plan de

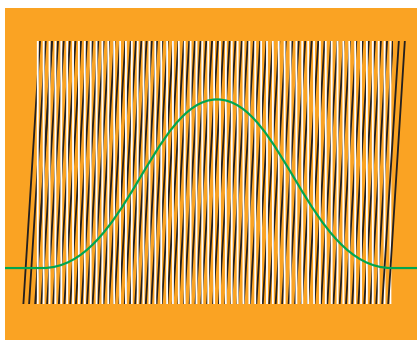
projection (a, au centre) se projettent en une série de droites parallèles. On obtient des figures de moirés en superposant ces deux réseaux de projection (a, à droite). Selon l'inclinaison du plan par rapport au cône, on obtient les figures d'intersection d'un cône par un plan : des paraboles (b), des ellipses (c) et des hyperboles (d).

horizontal, les figures moirées deviennent des ellipses (1c). À l'inverse, quand l'écartement diminue – le plan se redresse – les figures de moirés deviennent des hyperboles (1d).

Grâce aux figures de moirés, on peut recréer une courbe de Gauss en superposant deux séries de lignes parallèles. La courbe « en cloche » est fonction des probabilités d'une variable aléatoire soumise à de multiples causes indépendantes. La première série est composée de lignes verticales équidistantes. Les espacements de la seconde représentent les probabilités de chaque valeur que peut prendre la variable. En d'autres termes, les espacements de la première grille sont constants, ceux de la seconde sont plus serrés au centre. Les figures de moirés obtenues suivent la courbe de Gauss (voir la figure 2).

Les physiciens utilisent aussi les figures de moirés. Par exemple, une de leurs propriétés est d'amplifier les différences entre deux séries périodiques. Quand deux séries de lignes parallèles sont rigoureusement identiques, leur superposition est parfaite. À l'inverse, une légère différence dans la régularité ou dans les espacements apparaît quand les deux séries sont superposées (voir la figure à côté du titre).

En 1874, le physicien britannique Lord Rayleigh (1842-1919) a proposé d'utiliser cette caractéristique des figures de moirés pour vérifier la qualité des réseaux de diffraction : il s'agit de plaques (à l'époque en verre) transparentes striées de rainures régulièrement espacées. Ils sont employés pour étu-



2. Grâce aux moirés, on visualise les solutions du problème de l'interférence entre deux fonctions périodiques. Les lignes verticales blanches sont régulièrement espacées, alors que les espacements des lignes noires suivent une courbe de Gauss qui apparaît quand les deux grilles sont superposées.

der la lumière. Ainsi, les imperfections du réseau testé sont immédiatement détectées lorsque ce dernier est placé sur un réseau original dont on est sûr de la qualité.

Poursuivons avec la lumière ! En cristallographie, les phénomènes de moiré rendent d'innombrables services. Par exemple, quand deux fines lames de cristaux identiques sont superposées, des figures de moirés apparaissent au microscope électronique. Cependant, lorsque les cristaux sont plus épais, les électrons du microscope n'ont pas suffisamment d'énergie pour traverser les cristaux trop décalés l'un par rapport à l'autre : ils ne passent que lorsque la superposition est parfaite.

Les motifs obtenus résultent des interférences créées lors du passage des élec-

trons à travers les réseaux atomiques. Ces figures de moirés fournissent de multiples renseignements : par exemple, l'amplification inhérente aux moirés que nous avons décrite met en évidence des anomalies cristallines d'une taille inférieure au diamètre d'un atome (de l'ordre de quelques dixièmes de milliardièmes de mètre), soit à une précision 100 à 1 000 fois supérieure à la résolution du microscope.

Les figures de moirés aident aussi les cristallographes à déterminer, à l'aide des rayons X, la structure atomique de certaines molécules, telles les protéines, quand celles-ci sont sous une forme cristalline. Par ailleurs, les moirés aident à vérifier la qualité d'une lentille. Celle-ci est placée entre deux grilles et deux lentilles « parfaites » qui corrigent le grossissement des motifs. Quand la lentille testée renferme des distorsions, les lignes des moirés sont courbes.

Enfin, grâce aux moirés, on peut étudier la vitesse de dissolution, par exemple d'un sucre, dans une solution placée entre deux grilles de lignes parallèles : à mesure que le sucre se dissout, l'indice de réfraction de la solution évolue et modifie les figures de moirés que l'on observe.

Vasarély était très au fait de ce que se faisait en science. Connaissait-il pour autant toutes les applications possibles des figures de moirés ?

V. VASARELY, Monographie, vol. 1, Editions Griffon, Neuchâtel, 1965.

Fondation Vasarély, site Internet www.fondationvasarely.com



3. En 1953, Victor Vasarély réalise les œuvres de la série Transparence. Elles sont composées, d'une part, d'un motif de lignes noires et blanches (ci-dessus, à gauche) peint sur un support opaque, d'autre part, du même motif (dans certaines œuvres, il

s'agit du motif en négatif) reproduit sur un support transparent que l'on superpose au premier. Les mouvements du support transparent, par exemple un décalage ou une rotation (au centre et à droite), créent des figures de moirés.



La plongée sous-marine

ROLAND LEHOUCO • JEAN-MICHEL COURTY

Qui veut évoluer sous l'eau doit respecter des règles strictes, dictées par la physique et par la physiologie.

Le masque se plaque sur le visage : les oreilles font mal. Toute personne qui a déjà plongé le sait : la pression ambiante augmente à mesure que l'on s'enfonce sous l'eau (d'environ une fois la pression atmosphérique tous les dix mètres). À 20 mètres de profondeur, elle est ainsi le triple de la pression atmosphérique (c'est-à-dire la pression qui règne à la surface de l'eau plus la pression due à la couche d'eau). Les tissus mous qui constituent notre organisme sont à la pression ambiante. Composés pour l'essentiel de matières solides et d'eau, ils sont peu compressibles et ne changent quasiment pas de volume au cours d'une plongée.

En revanche, le comportement de l'air contenu dans le système respiratoire est tout autre. Les gaz sont beaucoup plus compressibles que les liquides (1 000 fois plus que l'eau). Dès le milieu du XVII^e siècle, l'Irlandais Robert Boyle et le Français Edme Mariotte énoncèrent une loi pour décrire leur compressibilité : «Le volume d'un gaz est, à température constante, inversement proportionnel à la pression qu'il subit.» Ainsi, quand la pression double, le volume qu'occupe un gaz est divisé par deux. Ceci explique que dès dix mètres de profondeur, le masque écrase le visage. Pour s'affranchir de cet inconvénient, les plongeurs munis de bouteilles doivent apprendre à insuffler de l'air dans leur masque par le nez à mesure qu'ils descendent. Pour éviter les douleurs dans les oreilles dues à une trop forte pres-

sion sur les tympans, ils soufflent par les voies nasales en se bouchant le nez, de sorte que la pression augmente dans les trompes d'Eustache (de petits conduits qui assurent l'équilibre de la pression sur les tympans).

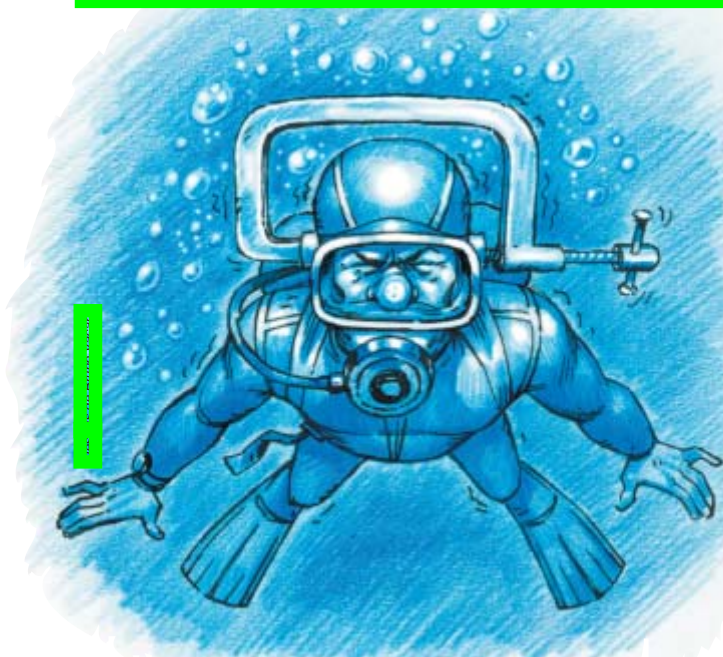
Les conséquences pratiques de la compressibilité de l'air sont encore plus nombreuses lorsque les plongeurs remontent. Après une expiration normale, la quantité d'air résiduelle dans les poumons est supérieure à trois litres, et elle n'est jamais inférieure à 1,5 litre, même lors d'une expiration profonde. Dans les poumons d'un plongeur qui remonte, cet air résiduel se dilate à mesure que diminue la pression ambiante. Si la remontée est trop rapide, le volume risque d'excéder la capacité maximale des poumons, qui avoisine six litres. Pour cette raison, les plongeurs qui remontent évitent de bloquer leur respiration ; ils expirent en permanence, de façon à éviter que les tissus qui composent leurs poumons ne se distendent jusqu'au point de rupture des membranes des cellules pulmonaires. Cet accident rare, mais redouté, est nommé la surpression pulmonaire.

Notons que la dilatation de l'air touche toutes les cavités intérieures, même celles dont on ne soupçonne pas l'existence ! Il arrive, par exemple, qu'un plombage mal réalisé contienne de minuscules cavités pleines d'air. En raison de la petite taille de l'ouverture qui les relie à l'extérieur, plusieurs minutes, voire plusieurs dizaines de minutes sont nécessaires pour que la pression qui règne à l'intérieur de l'une de ces petites bulles atteigne la pression ambiante. Toutefois, au bout d'une plongée de 40 minutes, par exemple, cet équilibre est atteint. Pendant la remontée, la pression dans les cavités n'a pas le temps de s'équilibrer. La surpression qui s'ensuit alors à l'intérieur des dents crée de violentes douleurs, peut faire éclater la dent, ou expulser le plombage. Si vous plongez, avez un bon dentiste !

DISSOLUTION DES GAZ

Les gaz ne se logent pas seulement dans les cavités, mais se dissolvent aussi dans les liquides. On trouve ainsi dans le sang tous les gaz que contient l'air. En 1803, le Britannique William Henry remarqua que la quantité de gaz qui se dissout dans un liquide est proportionnelle à la pression que ce gaz exerce sur le liquide, et que la quantité de gaz qui se dissout dans un liquide donné à une pression donnée dépend de la nature de ce gaz et de celle du liquide. De l'eau en contact avec de l'air contient des proportions d'azote et d'oxygène différentes de leurs proportions dans l'air.

Un litre d'eau à la pression atmosphérique contient ainsi quelque 15 milligrammes d'azote et 8 milligrammes d'oxygène, alors que sous la même pression, l'air contient quatre fois moins d'oxygène que d'azote. L'air qu'un plongeur respire sort des bouteilles, puis passe par le détendeur, qui le met à la pression ambiante. Puisque la pression augmente au cours de la descente, de plus en plus d'oxygène et d'azote se dissolvent



La pression augmente avec la profondeur ; elle écrase le nez et enfonce les tympans. Les plongeurs s'affranchissent des effets indésirables de la pression ambiante en soufflant dans leur masque.



ans le sang, jusqu'à une concentration d'équilibre. L'augmentation de la concentration en oxygène dissous dans le sang n'a pas de conséquences physiologiques néfastes, car notre métabolisme en consomme de grandes quantités. La quantité d'oxygène nécessaire est quelque 50 fois supérieure à la quantité totale d'oxygène dissoute dans le sang à pression atmosphérique. La majeure partie est transportée par l'hémoglobine.

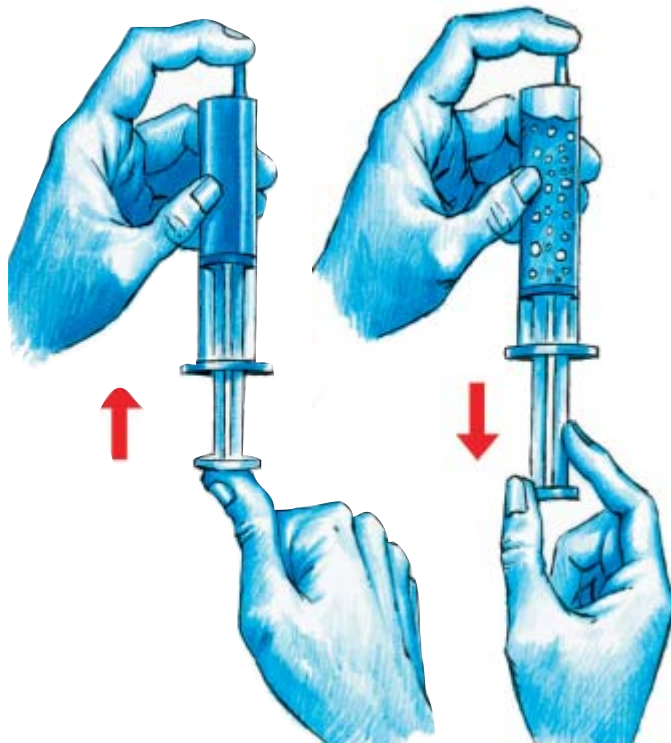
En revanche, l'augmentation de la concentration en azote dans le sang pose davantage de difficultés. Ce gaz se concentre dans le sang et dans tous les tissus issus de l'organisme, riches en eau. Quand une plongée dure longtemps, un équilibre s'établit entre la pression de l'azote dans les poumons et la quantité d'azote dissous dans les différents tissus. Lorsque le plongeur remonte, l'azote en excès repasse progressivement à l'état gazeux. Il est évacué petit à petit lors du passage du sang dans les poumons. L'évacuation est progressive et n'est possible que par les poumons : elle est longue. Lors des remontées trop rapides, l'azote dissous dans le sang et dans les tissus n'a pas le temps d'être éliminé. Des bulles de gaz restent piégées dans le sang et dans les tissus. Elles risquent de détériorer les tissus, créant des nécroses voire, plus grave encore, de rester bloquées dans des vaisseaux sanguins, empêchant le sang de s'écouler et provoquant des embolies. On peut imaginer les dégâts causés par ce type de décompression trop rapide en examinant le volume de gaz libéré par une bouteille de champagne que l'on sabre. Maintenu sous une pression de six atmosphères, le champagne (0,75 litre) contient environ 10 grammes de dioxyde de carbone et peut libérer plus de cinq litres de ce gaz dès qu'il est à pression ambiante.

LES RISQUES DE L'AZOTE

La procédure de remontée d'une plongée relativement brève à faible profondeur veut qu'un plongeur gravisse moins de 15 mètres par minute (la vitesse des petites bulles qui l'entourent). En revanche, une plongée longue ou profonde doit être suivie d'une ou de plusieurs pauses pendant la remontée : ce sont les paliers de décompression. Pendant un palier, le plongeur reste à la même hauteur pendant plusieurs minutes, ce qui laisse le temps à l'azote libéré dans le sang et dans les tissus d'être éliminé par les poumons. Le nombre de paliers nécessaires et leurs durées sont indiqués dans les tables de plongée en fonction de la durée du séjour sous l'eau et de sa profondeur.

Que se passe-t-il quand les circonstances imposent de ramener rapidement un plongeur à la surface? Pour éviter les risques de nécrose et d'embolie celui-ci doit impérativement être placé, le plus vite possible, dans un caisson hyperbare : dans un tel dispositif, on augmente artificiellement la pression de sorte que l'azote se redissout dans le sang et dans les organes. Ensuite, on fait des paliers de décompression progressifs, tels qu'ils auraient dû être.

Le volume occupé par un gaz est inversement proportionnel à la pression à laquelle il est soumis. Ainsi, le volume d'un ballon d'air est divisé par deux à dix mètres de profondeur, par trois à 20 mètres de profondeur et par quatre à 30 mètres. Un apnéiste qui, pour remonter, se fait tracter par un ballon, accélère au cours de son ascension, car le volume augmentant, la poussée d'Archimède, dirigée vers le haut, croît également.



Cette expérience, facile à réaliser avec de l'eau pétillante, montre que la solubilité d'un gaz dans un liquide diminue quand la pression baisse. Sous pression, l'eau gazeuse contenue dans la seringue dont on bouche l'extrémité (à gauche) ne libère aucune bulle. Quand on abaisse le piston (à droite), la pression à laquelle est soumise l'eau diminue : de nombreuses bulles de gaz apparaissent dans le liquide.

pratiqués dans la mer et qui permettent l'élimination progressive de l'azote.

Indépendamment de ces difficultés, l'azote reste dangereux. Au-delà d'une certaine concentration dans l'organisme, il devient toxique, et déclenche l'«ivresse des profondeurs». Les facultés physiques et intellectuelles s'engourdissent ce qui est très dangereux, car les réactions d'un plongeur soumis à cette forme de narcose sont imprévisibles et généralement inadaptées. Or, au-dessous de 40 mètres, tout plongeur est plus ou moins narcosé! Pour éviter ce danger, les plongeurs utilisent parfois, comme gaz de plongée du nitrox. C'est un mélange d'oxygène et d'azote contenant une proportion d'azote inférieure à celle de l'air. L'apport d'azote étant inférieur, le gaz redouté se dissout moins et plus lentement dans les tissus et dans le sang. En conséquence, moins de paliers de décompression sont nécessaires. Malheureusement l'oxygène lui-même devient aussi toxique au-delà d'une certaine pression!

Respirer de l'oxygène à plus de 1,6 atmosphère déclenche de graves malaises. L'oxygène entre pour 20 pour cent dans la composition de l'air. Pour un plongeur qui respire l'air contenu dans des bouteilles, l'oxygène devient toxique quand la pression ambiante atteint huit atmosphères, c'est-à-dire vers 70 mètres de profondeur. Qui veut descendre plus bas doit respirer un mélange à faible teneur en oxygène et en azote, tels l'héliox, un mélange d'hélium et d'oxygène, ou encore l'hydrox, composé d'hydrogène, d'hélium et d'oxygène.

Fascinante, l'exploration des fonds sous-marins nécessite des connaissances en physique et en physiologie que seuls délivrent les moniteurs compétents.

Roland LEHOUCQ est astrophysicien au Service d'astrophysique du CEA. Jean-Michel COURTY est physicien, chargé de recherche au CNRS. Philip FOSTER, *La plongée sous-marine à l'air*, Presses Universitaires de Grenoble, 1993. *La plongée sous-marine*, Encyclopædia Universalis.

Histoire des sciences

Lamarck, genèse et enjeux du transformisme 1770-1830

Pietro Corsi. CNRS Éditions, 2000
(424 pages, 290 francs)

A Paris, la statue de Jean-Baptiste de Monet, chevalier de Lamarck (1744-1829), accueille les visiteurs du Jardin des Plantes. Pour

tant, personne ne lui accorde le moindre regard : ni le promeneur qui profite du socle pour se reposer et admirer la perspective jusqu'à la Grande Galerie de l'évolution, ni les pompiers parisiens qui prennent appui contre son piédestal pour s'échauffer avant leur jogging matinal.

Lamarck est-il oublié? Pietro Corsi s'attaque d'abord à l'image qui en a été souvent donnée : celle de l'aristo-

crate-citoyen, pauvre professeur du Muséum, malheureux, incompris, méprisé de son vivant, ignoré après son décès. Il argumente : les travaux de Lamarck s'inscrivaient bien dans les préoccupations scientifiques de son époque, et ils ont été connus de ses contemporains, tant en France qu'à l'étranger, malgré l'hostilité du dogmatique Georges Cuvier (1769-1832). Il est abusif et simpliste de réduire la pensée lamarckienne à la loi de l'usage et du non-usage (elle énonce qu'un organe se développe s'il est utilisé et s'atrophie dans le cas contraire) et à l'hérédité des caractères acquis (un caractère développé pendant la vie d'un individu, se transmet aux générations suivantes).

Il est impossible, également, de considérer Lamarck comme un simple botaniste soutenu par Georges Louis Leclerc, comte de Buffon (1707-1788), ou un professeur de zoologie du Muséum, spécialiste des insectes et des vers (ou animaux sans vertèbres). En effet, le savant qui a donné à la biologie ses lettres de noblesse avait une pensée multiforme, qui incluait la physique (le

monde est régi par des lois), la chimie, la géologie, la météorologie (ses annuaires furent des succès populaires, car prévoir le temps est une préoccupation constante des hommes ; en revanche, ils furent jugés peu sérieux par les savants), la paléontologie, etc. Souvent, il s'est égaré, mais, comme l'écrit l'historien Goulven Lauren en 1997 : « Lamarck a eu autant raison que d'autres de son temps et autant que bien d'autres depuis. Et il a commis autant d'erreurs. »

P. Corsi s'est attaché à retracer l'évolution de la pensée du naturaliste depuis son adhésion au transformisme (ou transformation graduelle des formes de vie) dans son discours inaugural du cours de l'année 1800 puis dans les *Recherches sur l'organisation des corps vivants* (1802) et, enfin, dans la *Philosophie zoologique* (1809). La tâche était difficile : « Il est toutefois indéniable que les textes transformistes de Lamarck et ses œuvres en général sont difficiles à interpréter. Souvent, l'auteur ne se donne ni la peine d'indiquer clairement le type de problèmes qu'il aborde, ni de

distinguer explicitement les différents niveaux de son exposé. » Le style est aussi « répétitif et parfois maladroit ».

Malgré ces réserves, Lamarck est le fondateur du transformisme, à une époque où le fixisme était le concept dominant, défendu bec et ongles par Cuvier. Lamarck a imaginé l'échelle des êtres d'abord dans le sens d'une dégradation, des plus complexes aux plus simples, avant d'inverser son raisonnement. Il a expliqué le vivant par des forces matérialistes (c'était iconoclaste) s'est rallié aux générations spontanées, a introduit le temps en géologie, mais a nié les extinctions.

C'est cette pensée foisonnante, évolutive, contradictoire que P. Corsi décrypte. À la suite de Lamarck, des transformistes déterminés se sont déclarés, tels Étienne Geoffroy Saint-Hilaire (1772-1844) et Jean-Baptiste Bory de Saint-Vincent (1778-1846). En Angleterre, le débat a fait rage entre Robert Chambers (1802-1871) qui n'a pas osé signer son *Vestiges of the natural history of creation* (1844), parce qu'il défendait des idées évolutionnistes, et le géologue Charles Lyell (1797-1875) qui a beaucoup critiqué Lamarck dans son deuxième tome des *Principles of geology* (1832). Cette agitation éditoriale a préparé le terrain pour *De l'origine des espèces par voie de sélection naturelle* (1859) de Charles Darwin (1809-1882).

L'influence de Lamarck est attestée par la fréquentation de ses cours, de 1795 à 1823 ; ensuite, devenu aveugle, il a été remplacé par son suppléant, l'entomologiste André Latreille (1762-1833). Reprenant un travail initié par le professeur Max Vachon dans les années 1970, P. Corsi et R. Bange ont lancé une base de données recensant les élèves de Lamarck inscrits sur le registre des cours. Connus ou inconnus, ils sont plus de 100 étudiants en médecine ou en pharmacie, médecins, chirurgiens, naturalistes voyageurs ; leur origine géographique prouve le rayonnement du Muséum et de Lamarck sur les esprits éclairés de l'époque : Français, Allemands, Américains, Autrichiens, Belges, Brésiliens, Britanniques, Danois, Espagnols, Grecs, Hollandais, Roumains, Russes, Suisses. La recherche historique actuelle veut réunir le plus possible de renseignements sur ces auditeurs. Lamarck a maintenant son adresse internet : www.Lamarck.net! Qui peut encore parler de son isolement? La prévisior de la fille de Lamarck se trouve ainsi révisée : « La postérité vous admirera, elle vous vengera, mon père! »

Michèle FEBVRE
Université de Nice Sophia-Antipolis



Jean-Baptiste de Monet, chevalier de Lamarck, père du transformisme. Était-il méprisé de son vivant?

Marées, la vie secrète du littoral

Christophe Courteau.
Éditions Glénat, 2001
128 pages, 196 francs

Cet ouvrage est avant tout un recueil de superbes photographies. Les causes physiques responsables du phénomène des marées sont brièvement rappelés en introduction, mais que les principales conséquences physiologiques pour la flore et la faune. Par la suite, l'auteur envisage le « branle-bas de combat » déclenché par la marée descendante dans les populations animales littorales et les stratégies mises en place par les espèces vivantes : la fuite pour toujours vivre dans les eaux profondes pour les espèces mobiles, ou, pour les espèces fixées ou peu mobiles, la préparation à « quelques heures difficiles » liées au retrait de l'eau. Certains individus se retrouvent piégés dans des flaques d'eau, d'autres vivent en permanence dans le même endroit à marée haute et à marée basse et supporteront alors l'émersion. Puis, l'auteur appelle les « courbeurs de grèves » les espèces d'oiseaux littoraux plus ou moins spécialisés dans la prédation des invertébrés et des poissons, plus vulnérables à marée basse. Il termine par l'un des écosystèmes les plus intéressants en France pour ces aspects : où le phénomène de marée est si marqué, la baie du Mont-Saint-Michel. Le marnage peut y dépasser 16 mètres. Il évoque, grâce à de très belles images des moulins à marée tout près du Mont-Saint-Michel, mais pourquoi n'avoir pas parlé de l'usine marémotrice de la Rance si originale ?

Le texte est intéressant, riche en anecdotes, et les photographies sont de toute beauté, tant celles des milieux que celles de la faune sous-marine ou aérienne. Nous ne pouvons que recommander cet ouvrage sur un sujet traité de façon originale grâce à ces illustrations.

Nous cherchons souvent à aborder la physiologie de l'extrême par des espèces qui vivent dans les déserts ou sur les glaces de l'Antarctique, en bref, dans des milieux spectaculaires et exceptionnels, mais cet « extraordinaire » se situe pourtant souvent tout près de nous, et pas toujours où nous l'imaginons. Les espèces intertidales (qui vivent dans la zone de balancement des marées) sont confrontées à des situations extrêmes, des différences très importantes de



Le Lepadogaster de Gouan (poisson ventouse) se rencontre sous des pierres, dans les herminiers ou même dans des flaques entre 0 et 10 mètres de profondeur. Les nageoires pectorales et pectorales sont modifiées pour former une double ventouse ventrale. L'animal peut alors adhérer aux rochers et résister à la houle.

température, de salinité, d'éclairement de milieu (le passage de l'eau à l'air, et vice versa). Ces changements se produisent sur des périodes très courtes et à un rythme effréné (quatre fois par jour!). Par exemple, sur un estran vaseux, en été, une huître passe de 18 °C lorsqu'elle est sous six mètres d'eau à marée haute à 50 °C sur une vase en plein soleil à marée descendante. Les espèces animales incapables de contrôler leur température interne corporelle souffrent beaucoup de ces variations vertigineuses. Certaines espèces peuvent passer de l'eau salée à de l'eau quasi douce pour peu qu'un cours d'eau ou des pluies importantes dessalent le milieu à marée basse, alors gare aux chocs osmotiques ! Les différences de l'acidité de l'eau aussi sont parfois fortes, celle de l'eau de mer est très stable et basique, tandis que les rivières bretonnes qui s'écoulent sur les schistes sont parfois beaucoup plus acides. Comment survivent les espèces ? En fait, elles anticipent ces changements et se préparent, grâce à leurs coquilles étanches pour certains bivalves, à leur opercule pour les bigorneaux, ou encore en adhérant au rocher...)

L'un des grands dangers qui guette l'animal hors de l'eau au soleil est la déshydratation. La résistance de certains habitants de l'estran est extraordinaire. La variation de lumière, tant en longueur d'onde qu'en quantité, est parfois négligée par les scientifiques, mais pose pour autant certains problèmes. L'auteur présente une littorine (un mollusque gastéropode) qui a trouvé une stratégie d'animal terrestre, comme chez les escargots pulmonés, pour résister au passage à l'air

Chez certains poissons, la respiration à travers la peau est très efficace, grâce à de fins systèmes finement vascularisés à la périphérie de l'organisme. Le mucus produit par différentes cellules de la branchie ou de la peau joue alors un rôle déterminant. Finalement, beaucoup d'espèces de poissons littoraux parviennent à supporter l'émersion un certain temps : des expérimentations chez les poissons plats et les gobies, ont démontré des capacités exceptionnelles, largement supérieures à la durée d'un cycle de marée. La persistance d'un environnement humide est cependant indispensable. Les grands migrateurs (truite de mer et saumon, alose, anguille...) ne persistent pas longtemps dans cette zone de balancement des marées, mais ils « savent », à certaines étapes de leur migration (en descendant et en remontant) exploiter le phénomène pour ne pas passer trop brutalement d'un milieu à l'autre. Ils utilisent le mouvement des marées dans l'estran pour, peu à peu, sur quelques jours, s'adapter à la nouvelle salinité.

Toutes les espèces intertidales ont un rythme d'activité calqué sur le phénomène de marée : elles vivent à marée haute, respirent plus active, alimentation, excrétion, comportement reproducteur... et « survivent » à marée basse. Elles anticipent, et les travaux de chronobiologie démontrent de merveilleux rythmes internes liés à celui de la marée. Ainsi, après cinq ans en aquarium, en conditions constantes, un poulpe continuait encore à attendre la marée descendante.

Bien entendu, ces espèces ne sont pas dans des conditions idéales pour se développer et grandir vite : la même moule

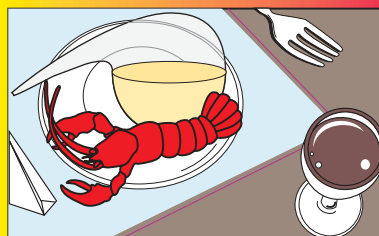
FRANCE INFO

POUR LA SCIENCE

vous invitent
à écouter la chronique
de Marie-Odile Monchicourt :

Info
Sciences

Tous les jours
sur France Info
à 15 heures 41,
17 heures 10,
20 heures 12,
22 heures 12,
et 23 heures 42



Les prochains
résultats présentés
dans la rubrique
Science et Gastronomie
seront annoncés sur
France Info
le 26 septembre 2001



placée en situation intertidale et en pleine eau au large ne montrera pas du tout les mêmes performances de croissance. Ces capacités à grandir plus faibles sont liées directement aux temps d'émergence qui interdisent aux animaux de se nourrir pendant la moitié du temps et aussi aux énormes coûts en énergie à dépenser pour résister aussi souvent à des conditions si difficiles.

Gilles BŒUF, Observatoire
océanologique de Banyuls-sur-Mer

Technique

Le Soleil, le génome et Internet. Les révolutions scientifiques et leurs instruments

Freeman Dyson.
Éditions Flammarion, 2001

Traduction de l'anglais

Une série de conférences que Freeman Dyson a données à New York et que Flammarion publie dans cet ouvrage constitue un bon

exemplaire de la diversité des propositions qu'un grand physicien peut tenir quand il sort de son laboratoire. Ces propositions paraîtront évidemment les plus pertinentes sur la pratique scientifique. Comme l'historien des sciences P. Galison, F. Dyson considère que l'évolution des instruments scientifiques est centrale pour comprendre la production des nouvelles connaissances. Il estime notamment que le scientifique doit concevoir lui-même ses instruments comme le fait l'astronome. Le choix des outils et des procédures de travail est très lié à la question du coût de la recherche. Le scientifique, s'il veut inscrire son travail dans la durée, doit chercher à minimiser les dépenses qu'il engage. Il vaut mieux, à l'instar des astronomes qui dressent la carte du ciel à l'aide de caméras électroniques, chercher des solutions simples, développées ensuite pendant de longues années, que de bâtir des projets spectaculaires sans lendemain parce que trop coûteux, comme Apollo ou la carte du génome. Le scientifique doit contrôler son budget, mais il doit aussi savoir se projeter dans l'avenir. Même si F. Dyson sait que l'anticipation est un genre difficile où l'on risque d'être démenti par les faits, il estime néanmoins que l'on a besoin de modèles

simples pour avancer et que, dans ce domaine, la précision et la cohérence sont plus importantes que la justesse. Contrairement aux romanciers d'anticipation, comme Albert Robida qui imaginait la technique en rêvant à ses usages, il part des grands domaines de la recherche scientifico-technique et de leur articulation. Ainsi, le génie génétique devrait permettre de créer des variétés de végétaux capables de convertir la lumière solaire en combustible. Si l'on ajoute au génie génétique une énergie solaire bon marché et Internet, de nouveaux modes de vie deviennent possibles : «Le plus modeste des villages mexicains offrira le même niveau de vie que Princeton.» En matière de génétique, il envisage également des perspectives plus inquiétantes, notamment celle de la «reprogénétique». En retirant certains gènes et en introduisant d'autres dans l'œuf fécondé, on pourrait développer de façon systématique une politique eugéniste et, à terme, créer des espèces humaines qui se sépareraient complètement les unes des autres.

A longue échéance, la seule façon d'éviter des conflits très graves entre ces espèces serait d'amener les plus aventureux à quitter la Terre et à émigrer. En effet, «une seule planète n'est pas suffisante pour expérimenter toutes les virtualités physiques et intellectuelles que renferme notre génome emporté par le grand élan de l'évolution». Cette prospective scientifico-technique rencontre ainsi, de temps à autre, le social. Et là, l'approche de F. Dyson reste celle d'un ingénieur. Il prône, par exemple, une politique égalitaire d'accès à Internet, parce que c'est «techniquement plus facile» à organiser que l'égalité dans l'accès au logement ou à la santé, comme si les barrières sociales étaient avant tout techniques! De même quand il se lance dans la rétrospective, il oublie toute la complexité du social. Après quelques pages de souvenirs familiaux où il nous explique que sa mère et ses tantes ont pu travailler parce qu'elles avaient des domestiques, il conclut, pour la génération de l'après-guerre, que «l'apparition des appareils électriques libéra les domestiques et enchaîna leurs patronnes». Est-ce une conséquence attristante du progrès technique ou, au contraire, faut-il considérer que l'organisation du travail domestique renvoie à une question spécifiquement sociale, la répartition des rôles féminin et masculin? Dès qu'il quitte le domaine de la techno-science, notre physicien manque manifestement des instruments adéquats pour penser le social.

Patrice FLICHY, Laboratoire
techniques, territoires et société
de l'Université de Marne-la-Vallée

Introduction à la psychologie du sport

Karine Bui-Xuan. Éditions Chiron
2000 (160 pages, 125 francs)

Des débuts du mouvement sportif en Grande-Bretagne, au XIX^e siècle, aux compétitions internationales modernes, le sport a subi

Elles ont conduit tous les acteurs sportifs (les cadres, les entraîneurs, les athlètes, les dirigeants, etc.), engagés dans une logique de performance, à solliciter les disciplines scientifiques, pour inventer et étudier les facteurs qui sous-tendent la réussite sportive. Le monde sportif a glissé d'une approche empirique de l'entraînement à une approche scientifique. En retour, les sciences ont accélééré cette transformation.

La demande du monde sportif n'est pas adressée d'emblée à tous les champs scientifiques, mais d'abord aux mathématiques, à la physique, à l'anatomie, à la physiologie, etc., qui décrivent les performances à partir de critères quantifiables. Plus récemment, la demande s'est orientée vers la psychologie. Pourquoi si tard? La psychologie du sport, issue de la psychologie générale et de la médecine sportive, a dû se libérer peu à peu de la tutelle de ces deux sciences pour traiter des questions spécifiques au sport. Cette séparation fut précoce dans les pays anglo-saxons ou de l'Europe de l'Est (fin des années 1960), mais plus tardive en France : au début des années 1980, elle profita de la création de structures d'enseignement et de recherche universitaires ou liées au ministère de la Jeunesse et des Sports.

Malgré ces structures, la psychologie du sport tarde à s'implanter dans le monde sportif, car elle est confrontée à des *a priori* réducteurs. Les dirigeants, les entraîneurs, les sportifs ne l'interrogent que lorsqu'ils sont confrontés à des dysfonctionnements, des échecs inexplicables ou des contre-performances. Cette discipline ne leur semble utile que dans la remédiation ou la réparation. Or, les connaissances développées en psychologie du sport devraient être intégrées dans la formation et dans l'entraînement des sportifs pour optimiser leurs performances. L'athlète se prépare sur le plan physique, technique, tactique. Et pour quoi pas mental? Il développerait ses



En France, la psychologie du sport n'est le plus souvent utilisée que pour comprendre les contre-performances des sportifs, alors qu'elle fait partie du programme d'entraînement des sportifs dans les pays anglo-saxons.

habiletés mentales pour exprimer son potentiel en compétition. Il ne serait plus question de réparer, mais de prévenir les contre-performances.

Enfin, la psychologie du sport a dû mal à sortir des débats théoriques qu'elle agite. Au niveau international, elle se construit sur une approche cognitive-comportementale. En revanche, en France, les approches clinique et cognitive-comportementale cohabitent. Comme ces deux courants théoriques ne s'accordent pas sur les conditions d'accès aux performances, l'un privilégiant le « lâcher prise », l'autre la mobilisation des habiletés mentales, ils ne conçoivent pas de manière identique la préparation psychologique des athlètes.

L'ouvrage de Karine Bui-Xuan ne rejette pas ces différences ; en ce sens, il paraît consensuel. Toutefois, au fil des pages, l'auteur accorde une place de plus en plus grande à la psychologie du sport clinique (l'expression est enfin employée par l'auteur à la fin du livre), soit une vision particulière de la psychologie du sport propre aux psychologues cliniciens. Ce livre ne traite donc pas de la psychologie du sport dans son ensemble, plutôt d'une psychologie du sport clinique. Il présente de façon complète les conceptions des psychologues cliniciens et répond au souhait de l'auteur de « rassembler des idées qui sont généralement exposées dans des

revues par l'intermédiaire d'articles » mais il n'approfondit pas les autres approches théoriques de la psychologie du sport. Les connaissances liées à l'approche cognitive-comportementale sont anciennes et présentées de manière superficielle. Elles ignorent la littérature anglophone, pourtant riche dans ce domaine. Cette disparité de traitement risque d'induire chez le lecteur une vision simpliste de l'approche cognitive-comportementale, voire la décrédibiliser. Au contraire de l'approche clinique, qui apparaît alors comme un recours incontournable. Ce décalage est regrettable et ne permettra pas au lecteur de se faire une idée complète de « ce qu'est vraiment la préparation psychologique ». Contrairement à ce qu'affirme l'auteur. Au début de l'ouvrage, deux définitions supplémentaires sur ce que l'auteur entend par sport et psychologie du sport auraient sans doute permis de mieux informer le lecteur du contenu.

L'ouvrage contribue à enrichir le regard du lecteur sur l'approche clinique de la psychologie du sport, mais il n'informe pas sur les connaissances spécifiques à cette discipline issues de l'approche cognitive-comportementale, pourtant les plus établies au niveau mondial et porteuses de nombreuses stratégies d'optimisation de la performance.

Jean-Philippe HEUZÉ, UFR STAPS de l'Université de Reims

SCIENCE

Une revue de la science et de la culture

www.pourlasience.fr

Commissariat général à l'égalité des territoires

1555-75278

POUR LA SCIENCE Directeur de la rédaction : Philippe Boulanger. Françoise Cinotti (Rédactrice en chef), Philippe Pajot, Loïc Mangin, Claire Le Poullennec, François Savatier, Laurence Plévert, Marie-Neige Corlonnier, Xavier Muller, René Cuillierier (Rédacteurs). Secrétariat de rédaction/Maquette : Annie Tacqueret, Céline Lapert, Sylvie Sobelman, Pauline Bilbault, Pierre Saminadin. Site Internet : Cyril Lamotte. Marketing et Publicité : Stéphane Montouchet, assisté de Séverine Merviel et Christophe Vitiello. Direction financière : Anne Gusdorf. Direction du personnel : Jean-Benoît Coutury. Fabrication : Jérôme Jalabert, assisté de Géraldine Le Coz. Presse et communication : Florys Grimaud, assistée de Isabelle Huet. Directeur de la publication et Gérant : Olivier Brossollet. Conseiller scientifique : Hervé This.

Ont également collaboré à ce numéro : Anne Briais, Jean Courtin, Jean-Michel Courty, Yves Fouquet, Jean-Claude Gérard, Évelyne Host-Platret, Christian Jeanpougin, Marie-Thérèse Landousy, Philippe Lejeune, Valérie Martin-Rolland, Michel Segonzac, Pascale Thiollier-Dumartin.

SCIENTIFIC AMERICAN Editor : John Rennie. Board of editors : Michelle Press, Mark Alpert, Carol Ezzell, Wyt Gibbs, Kristin Leutwyler, Madhusree Mukerjee, George Musser, Sasha Nemecek, Ricki Rusting, Sarah Simpson, Gary Stix, Philip Yam, Glenn Zorpette. Chairman Emeritus : John Hanley. Chairman : Rolf Gruebach. President and Chief Executive Officer : Gretchen Teichgraber. Vice-President : Chuck McCullagh and Frances Newburn.

PUBLICITE France

Chef de Publicité : Jean-François Guillotin (j.f.guillotin@pourlasience.fr), assisté d'Irène Limon

8 rue Férou 75278 Paris Cedex 06

Tél. : 01 55 42 84 28 ou 01 55 42 84 97

Télécopieur : 01 43 25 18 29

Étranger : 415 Madison Avenue, New York

N.Y. 10017 - Tél. (212) 754.02.62

SERVICE ABONNEMENTS

Anne-Claire Temois : 01 55 42 84 04

SERVICE DE VENTE RESEAU NMPC

Christophe Vitiello : 01 55 42 84 81

DIFFUSION DE LA BIBLIOTHEQUE

POUR LA SCIENCE

Canada : Edipresse : 945, avenue Beaumont, Montréal, Québec, H3N 1W3 Canada.

Suisse : Servidis : Chemin des châlets, 1979 Châtaignes - 2 - Bogis

Belgique : La Caravelle : 303, rue du Pré-aux-oies, 130 Bruxelles.

Autres pays : Éditions Belin : 8, rue Férou - 75278 Paris Cedex 06.

Toutes demandes d'autorisation de reproduire, pour le public français ou francophone, les textes, les photos, les dessins ou les documents contenus dans la revue «Pour la Science», dans la revue «Scientific American», dans les livres édités par «Pour la Science» doivent être adressés par écrit à «Pour la Science S.A.R.L.», 8, rue Férou 75278 Paris Cedex 06.

© Pour la Science S.A.R.L.

Tous droits de reproduction, de traduction, d'adaptation et de représentation réservés pour tous les pays. La marque et le nom commercial «Scientific American» sont la propriété de Scientific American, Inc. Licence accordée à «Pour la Science S.A.R.L.»

En application de la loi du 11 mars 1957, il est interdit de reproduire intégralement ou partiellement la présente revue sans autorisation de l'éditeur ou du Centre français de l'exploitation du droit de copie (20, rue des Grands-Augustins - 75006 Paris).

Arts et sciences

Rencontres entre artistes et mathématiciennes. Toutes un peu les autres

Thérèse Chotteau, Francine Delmer, Pascale Jakubowski, Sylvie Paycha, Jeanne Peiffer, Yvette Perrin, Véronique Roca, Bernadette Taquet. Éditions L'Harmattan, Bibliothèque du féminisme, 2001.

164 pages, un CD, 120 francs)

Arts et sciences. Sciences et femmes. Femmes et arts. Trois thèmes propices à des débats variés. Exemples. Pourquoi les

mathématiciennes ont-elles pris de mal à percevoir dans certains domaines, telles les mathématiques et la musique, que dans d'autres? Quel statut donner aux cires anatomiques du XVIII^e siècle ou aux modèles géométriques en plâtre des années 1900 : devons-nous les considérer comme des objets scientifiques parce qu'ils étaient à leur création, comme des objets d'art ou encore comme de pauvres témoins d'un passé mort? Y a-t-il un sens à analyser mathématiquement la géométrie sous-jacente à une statue? Sur ces questions, et sur d'autres, les «auteures», deux sculptrices, une musicienne et cinq mathématiciennes, proposent «une confidence qui voudrait laisser percer l'essentiel».

Sans féminisme militant, elles montrent, simplement, comment être femme et artiste, femme et mathématicienne. Elles cherchent non pas à confronter des idées, mais à exprimer leur vécu. Accompany d'un CD avec une œuvre électroacoustique de Pascale Jakubowski, illustrée par de nombreuses photographies de statues et de surfaces géométriques, le livre offre un aperçu sur le travail des artistes. Il ne défend pas de point de vue général que l'on puisse résumer et discuter. Rares d'ailleurs y sont les passages collectifs ; la plupart sont signés de l'une des huit auteures. Voici trois des points qui m'ont intéressé.

Deux autoportraits de femme, en cire. L'un date du XVIII^e siècle. D'après les photographies qui en sont données, on le verrait criant de vérité. L'autre, dû à Véronique Roca et exécuté en l'an 2000, ne cherche pas, au contraire, à être ressemblant. Si l'artiste du XVIII^e siècle a voulu immortaliser ses traits, celle du XX^e souligne l'éphémère par un autoportrait composé d'une répétition d'éléments de son corps pris dans des états différents. Aujourd'hui, où une simple photogra-

phie suffit à attraper la ressemblance, les artistes aspirent à voir, au-delà, ce qui est derrière elle.

De son côté, Jeanne Peiffer raconte le piège où elle s'est trouvée. Née dans un pays (le Luxembourg) où les filles entamaient rarement des études de sciences, elle a choisi les mathématiques par défi. Croyant échapper ainsi au piège que la société lui tendait, elle s'est en réalité enfermée de plus belle. Bien plus tard, en effet, elle s'est rendu compte que les mathématiques, à vrai dire, l'intéressaient peu. Où aller, après tant d'années consacrées à les apprendre? Pour ne pas tout perdre, elle s'est faite historienne des mathématiques. Selon elle, les femmes sont rares en mathématiques à cause de la forte structuration sociale de ce milieu.

Enfin, j'ai apprécié le tranquille irrespect de Sylvie Paycha. Jeune mère, perdue entre la course entre les enfants qui rient et le lait qui va brûler dans la casserole, cette mathématicienne n'a pas toujours le temps d'approfondir les idées qui lui viennent. Comment s'y prendre, alors? Comme avec les marchandises qu'elle empile en vitesse dans son caddy juste avant que le supermarché ne ferme. «Choisir une idée, sans avoir le temps de la soupeser, ni de regarder en arrière pour voir ce qu'on aurait oublié. On la jette dans le panier, où elle se mélange avec les autres. Peut-être sera-t-elle déjà avariée demain ou au mieux deséchée. Au moins elle n'aura pas brûlé.»

Didier NORDON



Autoportrait (détail) (cire), de Véronique Roca