

Petit trou noir

Depuis environ 10 ans, les astronomes se sont persuadés que le centre des grandes galaxies – dont la Voie lactée – renferment chacun un énorme trou noir de plusieurs millions, voire des milliards de masses solaires. La matière qui se trouve violemment aspirée par ces monstres est chauffée à des températures considérables et émet de grandes quantités de rayonnement, ce qui expliquerait que les noyaux des galaxies produisent beaucoup d'énergie à partir d'un très faible volume (du même ordre que celui du Système solaire pour une production d'énergie équivalente à celle de dizaines de milliards de soleils). Cependant, en étudiant la vitesse de révolution des étoiles de la galaxie M33 autour du centre de cette galaxie, David Merritt, de l'Université d'État du New Jersey, et ses collègues viennent d'établir qu'elle renferme un trou noir central n'excédant pas (si toutefois il existe) 3 000 masses solaires (1 000 fois moins que le plus petit trou noir central connu). Si ce résultat était confirmé, il remettrait en cause un des rares faits qui semblaient acquis quant au mécanisme encore inconnu de formation des galaxies.

La croissance des dinosaures

Une équipe de paléontologues franco-portugaise a montré que les mécanismes de croissance des os des théropodes, des dinosaures carnivores, sont plus proches de ceux des oiseaux, leurs descendants, que de ceux des tortues et des crocodiles. Ils ont étudié les os d'embryons vieux de 150 millions d'années, provenant du gisement de Lourinha, à 60 kilomètres de Lisbonne. Jusqu'à présent, la biologie des théropodes était mal connue, car les fossiles découverts étaient surtout des embryons de dinosaures végétariens du Crétacé supérieur, vieux de 70 millions d'années.

Les moustaches du phoque

Le phoque suit une proie en détectant, grâce à ses moustaches, le sillage du poisson. Une équipe allemande a simulé le mouvement d'un poisson au moyen d'un



Science

sous-marin miniature et a constaté que le phoque est capable de suivre l'engin à longue distance, même dans l'obscurité, alors qu'il échoue lorsque ses moustaches sont recouvertes par un bandeau. Les moustaches, très sensibles aux vibrations de l'eau, joueraient le rôle d'un détecteur. C'est la première fois que l'on identifie chez un animal ce type de repérage des proies.

Prélude aux éruptions

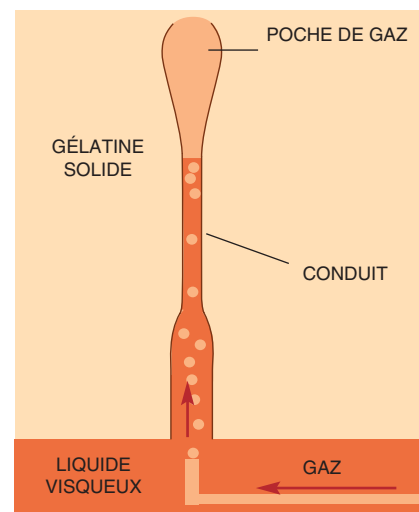
Un modèle expérimental de volcan explique l'origine des jets de gaz précurseurs des éruptions volcaniques.

Le 15 juin 1991, aux Philippines, le volcan Pinatubo se réveille après 500 ans d'inactivité, envoyant dans l'atmosphère un immense nuage de gaz et déversant sur ses flancs une lave qui balaiera tout sur son passage faisant 850 morts et détruisant les habitations d'un million de personnes. Le bilan a été lourd, mais il a pourtant été limité par la prévision de la catastrophe. Pendant les deux mois qui l'ont précédée, des éruptions annón-ciatrices riches en gaz et de faible volume avaient alerté les autorités du réveil du volcan et des mesures avaient été prises pour évacuer les personnes vivant à proximité. Ces jets précurseurs ont aussi été observés pour d'autres volcans, tels la Soufrière de l'île de Montserrat, aux Antilles, et le volcan Usu, au Japon. Pourquoi ces gaz initialement dissous dans le magma précédent-ils le magma? Empruntent-ils un réseau connecté de fissures? Thierry Menand et Stephen Tait, de l'Institut de physique du Globe de Paris, proposent un modèle expérimental de volcan, en laboratoire, qui rend compte des caractéristiques des jets de gaz précurseurs.

La source d'un volcan est une chambre magmatique, réservoir situé à quelques kilomètres sous la surface, où le magma s'accumule jusqu'à ce que la pression qu'il subit l'en éjecte. À mesure que le magma remonte, sa pression diminue et les gaz auparavant dissous s'échappent de la lave. L'observation des fissures magmatiques fossiles montre que ces gaz ont tendance à s'accumuler à l'avant des fissures et à former une poche. Dans le modèle expérimental de T. Menand et S. Tait, un cube de gélatine de 50 centimètres de hauteur simule la roche, et un liquide visqueux le magma. Le cube de gélatine obstrue hermétiquement la cuve contenant le liquide, et l'on en augmente la pression. On amorce un conduit dans la gélatine par une encoche. À mesure que l'on augmente la pression, le liquide se fraye un chemin vertical dans la gélatine. Puis, on stabilise la pression, et le liquide continue à monter lentement; on injecte alors dans la cuve quelques bulles d'air comprimé qui déforment la fissure: une poche de gaz apparaît à l'avant du liquide alors que celui-ci continue de progresser vers le haut. En même temps qu'elle se forme, la poche s'élève de plus en plus vite dans la gélatine. Le liquide, à cause

des frottements avec les parois du conduit, monte moins vite. Malgré l'accélération, le débit du liquide reste constant, ce qui signifie que la section du conduit diminue: la cheminée se referme derrière la bulle de gaz, laquelle continue à progresser et perce la surface de la gélatine quelques instants avant le liquide.

Ce modèle est-il une bonne image de la réalité? La pression des gaz précurseurs augmente jusqu'à la limite de rupture de la gélatine. Le modèle montre qu'une fois cette limite atteinte, la poche de gaz se désolidarise du magma, car, du fait de sa faible viscosité, elle avance plus vite que le liquide. Dans ce modèle, le liquide emprunte le même chemin que les gaz et l'intervalle de temps entre le jet précurseur et la véritable éruption dépend de la viscosité du liquide. Dans la réalité, cet intervalle dépend de la nature des gaz dissous et des pressions auxquelles ils quittent le magma. Le modèle prévoit toutefois qu'il est plus court dans le cas des volcans à lave basaltique, peu visqueuse, que dans les volcans où la lave contient de la silice. Cette hypothèse semble confirmée par le volcan Parícutin, au Mexique, où des gaz et des cendres sont émis quelques heures avant que le magma n'atteigne la surface, alors que dans le cas des éruptions à silice, des semaines, voire des mois (pour le volcan Pinatubo), séparent les deux événements.



Dans ce modèle, une couche de gélatine simule la roche et un liquide visqueux le magma. Les gaz injectés s'accumulent en une poche qui progresse dans la gélatine, devant le «magma».

AP Photo/Koji Sasahara



Un scribe moléculaire

Grâce à des mutations et à des sélections, des biologistes ont mis au point une molécule d'ARN qui copie des molécules... d'ARN.

La biologie des êtres vivants est régie par un dogme : l'ADN code l'ARN et l'ARN code les protéines (l'ADN est l'acide désoxyribonucléique, l'ARN l'acide ribonucléique). En a-t-il toujours été ainsi ? De nombreux biologistes pensent que cette organisation requiert trop d'acteurs moléculaires pour être apparue *ex nihilo*. Selon eux, la vie serait née avec l'apparition de molécules simples capables de se répliquer elles-mêmes sans avoir recours à une machinerie complexe : le meilleur candidat pour ces molécules primordiales est l'ARN, dont certains, nommés ribozymes, sont doués de propriétés catalytiques, c'est-à-dire qu'ils favorisent certaines réactions chimiques. Toutefois, aucun fossile moléculaire d'un tel ARN n'existe et aucun ribozyme connu ne copie des brins d'ARN. Pourtant, l'équipe de David Bartel, de l'Institut de technologie du Massachusetts, a mis au point une molécule d'ARN qui copie un brin modèle.

Les ARN sont constitués de nucléotides, c'est-à-dire d'assemblages d'une molécule azotée, nommée base, choisie parmi quatre (l'uracile, la cytosine, la guanine et l'adénine), d'un sucre (le ribose) et d'un acide phosphorique. C'est l'enchaînement des bases qui code l'information contenue dans l'ARN. Les bases sont complémentaires deux à deux : ainsi, la cytosine se lie à la guanine et l'adénine à l'uracile. La copie d'un brin d'ARN se fait par ajout successif des nucléotides complémentaires du brin modèle. Les nucléotides complémentaires sont présents sous une forme triphosphate : lorsqu'ils sont liés, une molécule de pyrophosphate (deux molécules d'acide phosphorique) est libérée. La copie démarre quand, sur le brin modèle, est déjà fixé un petit brin complémentaire initiateur.

Les biologistes avaient à leur dispo-

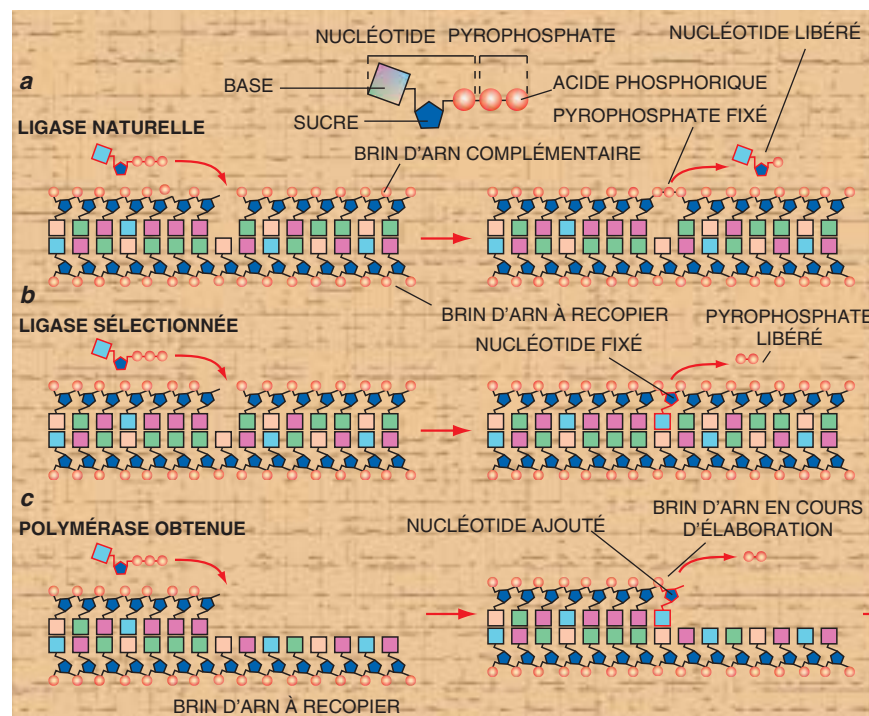
sition une molécule d'ARN catalytique connue, une ligase, qui lie deux petits brins d'ARN, chacun complémentaire d'une partie d'un brin modèle plus grand (voir la figure). Cependant, cet ARN fixait le pyrophosphate entre les deux brins complémentaires et libérait le nucléotide, alors qu'elle aurait dû faire exactement l'inverse ! L'un des deux brins complémentaires joue le rôle d'initiateur. Ainsi, les biologistes ont-ils tenté de modifier la ligase afin qu'elle libère le pyrophosphate et fixe le nucléotide et afin qu'elle lie à l'initiateur non pas un brin, mais un seul nucléotide. En fait, ils cherchaient à transformer cette ligase en une polymérase, le ribozyme qui copie les brins d'ARN nucléotide par nucléotide.

En 1993, les biologistes avaient franchi la première étape : la ligase qu'ils ont obtenue à partir de séquences aléatoires sélectionnées libérait bien le pyrophosphate et fixait le nucléotide, mais elle restait une ligase. Aujourd'hui, ils ont soumis le noyau catalytique – la partie active

de la molécule – de la ligase de deuxième génération à des mutations aléatoires. Puis, ils ont modifié certaines parties pour obtenir les propriétés recherchées : par exemple, ils ont rajouté un fragment aléatoire de 76 nucléotides à une extrémité de la molécule originale afin de créer une cavité où s'insère l'ARN à copier. Auparavant, l'activité de la ligase dépendait de la séquence de l'ARN à copier. La nouvelle molécule (qui agit maintenant comme une polymérase) recopie tous les brins d'ARN quelle que soit leur séquence.

Les biologistes ont sélectionné les différents ARN obtenus et étudié leurs propriétés de répllication. Après 18 opérations de mutation et de sélection, ils sont parvenus à un ARN de 189 nucléotides capable de copier jusqu'à 14 nucléotides de n'importe quel brin modèle, avec une fidélité de 96,7 pour cent.

La vie est-elle née avec les ARN ? Les travaux de D. Bartel rendent cette hypothèse plausible.



Les ARN catalytiques, ou ribozymes, sont des ARN qui favorisent une réaction. Un ribozyme naturel, nommé ligase, lie deux brins d'ARN complémentaires sur un brin modèle avec une molécule de pyrophosphate et libère un nucléotide (a). À partir de cette ligase naturelle, des biologistes ont d'abord sélectionné une ligase (b) qui lie deux brins complémentaires par un nucléotide. Puis, dans un deuxième temps, ils ont mis au point une polymérase qui, à partir d'un brin initiateur copie un brin modèle par additions successives de nucléotides (c).

Vous cherchez un article scientifique ou technique ?

//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

Des réseaux de neurones auto-organisés

Des robots commandés par un nouveau type de réseau de neurones s'adaptent à des situations nouvelles, sans oublier ce qu'ils ont déjà «appris».

Les réseaux de neurones artificiels sont de plus en plus utilisés pour contrôler l'activité de robots que l'on nomme animats. Au cours de l'évolution d'un animat dans son environnement, son réseau de neurones change de configuration en fonction des stimulus qu'il reçoit, mais aussi de ses configurations antérieures. Ainsi, le réseau de neurones contient, à chaque instant, des informations sur les objets qui entourent l'animat et sur les conséquences de ses actions passées. Contrairement aux systèmes plus rigides dont le comportement est déterminé par des décisions définies à l'avance, l'animat est doué d'une certaine autonomie et d'une capacité d'apprentissage. Toutefois, les réseaux de neurones artificiels ont un grave défaut : ils ont tendance à apprendre de façon globale (ils codent toutes les informations avec le même niveau de priorité) et une modification notable de l'environnement peut les conduire à remettre en cause l'ensemble des comportements appris indépendamment du contexte : c'est l'«oubli catastrophique». Nous avons proposé et testé, au cours de simulations informatiques, un nouveau modèle de réseau capable de s'en affranchir.

Le plus souvent, pour éviter l'oubli catastrophique, on hiérarchise à l'avance les informations traitées par l'animat. Le

réseau de neurones est ainsi divisé en modules spécialisés, par exemple, dans les tâches de reconnaissance de l'environnement ou encore de gestion des mouvements. Toutefois, cette méthode est peu satisfaisante, parce que cette modularisation reflète davantage les *a priori* du chercheur qu'une division des tâches fondée sur la façon dont le robot extrait des informations pertinentes de son environnement.

Pour créer un nouveau modèle, nous nous sommes inspirés de la neurobiologie. En effet, selon les neurobiologistes, l'information traitée par un réseau de neurones – le cerveau – est codée non seulement dans les motifs formés par les groupes de cellules qui sont actives à chaque instant, mais aussi dans l'évolution au cours du temps de ces motifs. Cette dépendance temporelle provient, semble-t-il, du fait qu'un neurone n'est activé qu'après une phase pendant laquelle il ne fait qu'accumuler les impulsions émises par les neurones auxquels il est connecté. En conséquence, il change d'état d'autant plus vite qu'il a été fortement stimulé. Nos neurones artificiels simulent ce comportement de façon simple : à chaque séquence, le programme ne réévalue que l'état des neurones les plus stimulés à la séquence précédente. Nous avons appelé ce mode d'activation, le mode ordonnancé.

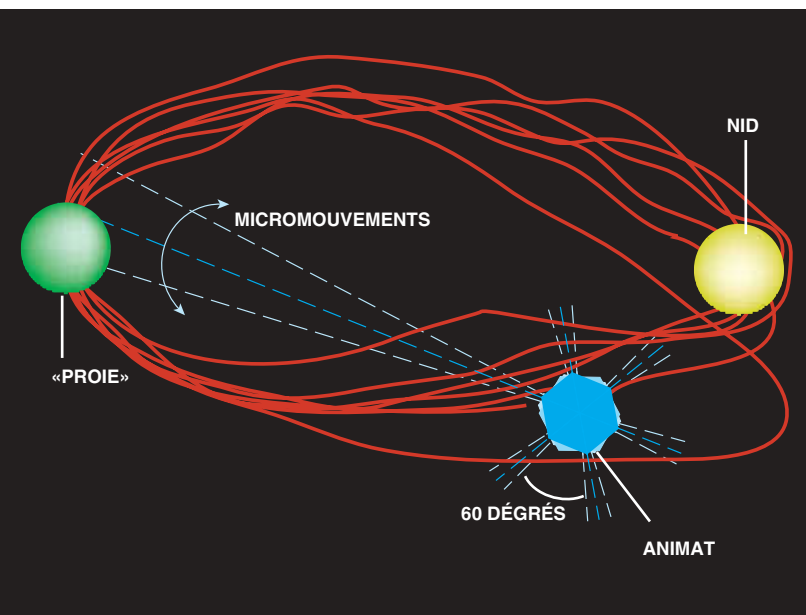
Notre animat évolue dans une pièce à deux dimensions contenant une «proie» qui se déplace aléatoirement et qu'il doit capturer et ramener à un «nid» fixe. Nous avons délibérément limité ses capacités sensorielles et motrices : il ne dispose d'aucune information sur sa propre position dans la pièce, ni sur celle du nid, ni sur la distance qui le sépare de sa cible. Ses six «yeux» lui indiquent simplement le secteur (de 60 degrés) où se trouve la cible et il ne perçoit les murs qu'à très faible distance. Par ailleurs, lors de sa «venue au monde», le réseau neuronal de l'animat est «brut» (nous n'avons introduit aucun module préalable) et toutes les cellules sont connectées entre elles. Au cours de l'évolution de l'animat, ces connexions sont renforcées ou affaiblies par un algorithme d'apprentissage qui favorise celles qui accompagnent les comportements à succès.

À cause du mode d'activation ordonnancé, tout se passe comme si l'influx nerveux se propageait à une vitesse finie dans le réseau. Ainsi, le nombre de neurones activés dans un comportement donné est limité par la vitesse à laquelle l'environnement évolue autour de l'animat. L'algorithme d'apprentissage ayant tendance à fixer les motifs de neurones qui provoquent les «bons» comportements, des populations de neurones différenciées émergent et l'animat les utilise pour réaliser des comportements différents : tout se passe comme si une structure en modules apparaissait spontanément dans le réseau.

Cette modularité n'est pas contenue à l'avance dans la structure du réseau. Si elle est difficile à observer directement, elle est mise en évidence par le fait, vérifié lors des simulations, que le phénomène d'oubli catastrophique disparaît. En outre, on observe d'intéressantes conséquences dans le comportement de l'animat. Il développe, notamment, des stratégies de repérage et d'approche de la cible fondées sur des micro-mouvements (de petites fluctuations du mouvement autour d'une trajectoire moyenne). Ces stratégies sont intéressantes parce qu'elles rappellent celles que l'on observe chez un grand nombre d'animaux et parce qu'elles émergent spontanément (nous n'avons en rien favorisé leur apparition).

Ce modèle offre un nouveau champ de recherche. Il devra être testé pour des tâches et des environnements dynamiques plus complexes. Ainsi, l'animat devra bientôt affronter le monde réel...

Guillaume BELSON et Hédi SOULA,
Institut national des sciences appliquées de Lyon



Le robot (en bleu) commandé par un réseau de neurones ne perçoit de ses cibles, la «proie» (en vert) ou le «nid» (en jaune) que leur présence ou non dans des secteurs de 60 degrés. Grâce au nouveau réseau de neurones, il adopte spontanément une stratégie «astucieuse» : à l'aide de petites fluctuations de sa trajectoire, il maintient sa cible entre deux capteurs ce qui améliore sa précision. Ces «micro-mouvements» rappellent, par exemple, ceux qui animent les yeux des mouches. Dans le cas où la cible est immobile, cette stratégie qui vise à maintenir une orientation moyenne constante entre le robot et la cible explique l'allure quasi elliptique de la trajectoire qu'il apprend, de lui-même, à parcourir.