

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.

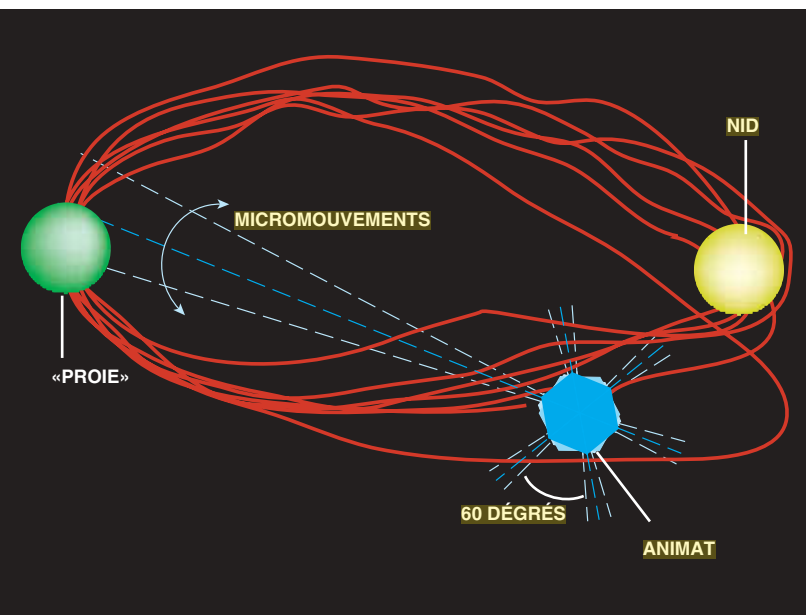
Notre animat évolue dans une pièce à deux dimensions contenant une «proie» qui se déplace aléatoirement et qu'il doit capturer et ramener à un «nid» fixe. Nous avons délibérément limité ses capacités sensorielles et motrices : il ne dispose d'aucune information sur sa propre position dans la pièce, ni 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.