

Chaque tentative de réforme devra combattre des habitudes fermement établies: la façon dont les statistiques sont enseignées dans les universités, celle dont les résultats des études sont exploités et interprétés et comment ils sont ensuite relayés dans les revues spécialisées. Mais au moins, de nombreux scientifiques ont admis l'existence d'un problème. Grâce à des chercheurs comme John Ioannidis, les préoccupations des statisticiens ne sont plus vues uniquement comme de la pure théorie.

**Il est
pratiquement
impossible
d'arracher les
auteurs à leur
sacro-sainte
valeur-*p***

Pour améliorer la situation, les statisticiens ont quelques outils à leur disposition. Par exemple, les chercheurs pourraient toujours publier également les tailles d'effet et les intervalles de confiance (*voir les Repères, page 6*) obtenus. Ces valeurs expriment ce que ne peut exprimer la valeur-*p* seule: la portée et l'importance relative de l'effet.

Beaucoup plaident aussi pour remplacer la valeur-*p* par des méthodes fondées sur la vision de Thomas Bayes, mathématicien britannique du XVIII^e siècle. Cette approche décrit une façon de penser les probabilités comme la plausibilité d'un résultat, plutôt que la fréquence potentielle de ce résultat. On introduit certes par là une subjectivité dans les statistiques, que les pionniers du début du XX^e siècle voulaient éviter à tout prix. Pourtant, la statistique bayésienne simplifie l'intégration des connaissances contextuelles sur le monde et permet de calculer comment les probabilités sont modifiées par des preuves nouvellement acquises.

D'autres soutiennent une approche plutôt pluraliste: les scientifiques devraient utiliser plusieurs méthodes sur un même ensemble de données. Lorsque les résultats différeront, les chercheurs devront être plus créatifs pour

découvrir à quoi cela est dû. La compréhension de la réalité sous-jacente y gagnerait.

Aux yeux de Simonsohn, la meilleure protection pour un chercheur est de tout montrer. Les auteurs devraient garantir leur travail «sans *p*-hacking» en explicitant le choix de la taille de leurs échantillons, les données éventuellement mises de côté ainsi que les manipulations mises en œuvre. Aujourd'hui, aucune de ces informations n'est disponible ni vérifiable.

DES ÉTUDES EN DEUX ACTES

À l'instar de la mention «Les auteurs n'ont pas d'intérêts financiers dépendants du contenu de cet article» encore courante aujourd'hui, ces informations feraient la différence entre faute scientifique involontaire ou délibérée. Lorsqu'une telle déclaration sera monnaie courante, le *p*-hacking sera éliminé. Ou au moins son absence sera-t-elle remarquée par le lecteur et lui permettra un jugement plus éclairé.

L'analyse en deux étapes ou «*preregistered replication*» (réplication préenregistrée) est une idée qui va dans ce sens, et elle gagne du terrain. Dans cette approche, les études exploratoires et confirmatoires sont abordées différemment et clairement identifiées. Au lieu, par exemple, de conduire quatre petites expériences et de réunir les résultats dans un article, les chercheurs balaiseraient d'abord le domaine pour récolter des observations intéressantes avec deux petites études exploratoires, sans trop se préoccuper des fausses alertes. C'est seulement après, sur la base de ces données, qu'ils concevraient une étude qui confirmera, peut-être, leurs résultats, et préenregistreraient publiquement leurs intentions dans une banque de données publique, telle l'Open Science Framework. Ils publieraient ensuite leurs résultats, accompagnés de ceux de l'étude exploratoire dans un article d'une revue habituelle. Une telle approche laisse beaucoup de liberté, explique le politologue et statisticien Andrew Gelman, de l'université Columbia à New York. Mais elle est aussi suffisamment rigoureuse pour diminuer le nombre de fausses découvertes.

Plus généralement, l'heure est venue pour les scientifiques de prendre conscience des limites des statistiques conventionnelles. Avant tout, une estimation scientifique sérieuse de la plausibilité des résultats devrait être effectuée dès leur analyse: quels résultats ont été apportés par des recherches similaires? Existe-t-il un mécanisme qui pourrait les expliquer? Les résultats se recoupent-ils avec l'expérience clinique? Voilà les questions décisives! En y répondant dans de prochains travaux, Motyl peut encore espérer accéder au succès qu'il espère! ■

BIBLIOGRAPHIE

D. BENJAMIN ET AL.,
Redefine statistical significance, 2017
<https://psyarxiv.com/mky9j>

B. NOSEK ET AL.,
Restructuring incentives and practices to promote truth over publishability,
Perspect. Psychol. Sci.,
vol. 7, pp. 615-631, 2012.

C. LAMBDIN
Significance tests as sorcery : Science is empirical – significance tests are not,
Theory & Psychology,
vol. 2, pp. 67-90, 2012.



Le *big data*, un flot de données
venant de toutes parts.



© Shutterstock.com/Dmitry Rybin

LES DÉFIS DES BIG DATA

Comment faire face à ce déluge de données qui constitue les *big data* ? Comment les stocker, les exploiter, les analyser, les valoriser... ? Les algorithmes et les calculateurs à qui ces tâches incombent doivent être à la hauteur. La recherche sur ces deux aspects essentiels du traitement des *big data* est particulièrement active. Le développement des intelligences artificielles, « éduquées » par des quantités énormes de données, est une illustration de ces efforts. Ces systèmes sont parfois si performants que certains y voient une menace sur la façon de faire de la science.

L'ESSENTIEL

● L'intelligence artificielle telle qu'on l'imaginait dans les années 1950 a été plus difficile à développer que prévu.

● Depuis quelques années, le domaine a connu un grand renouveau avec les techniques de l'apprentissage profond, inspirées des réseaux de neurones du cerveau.

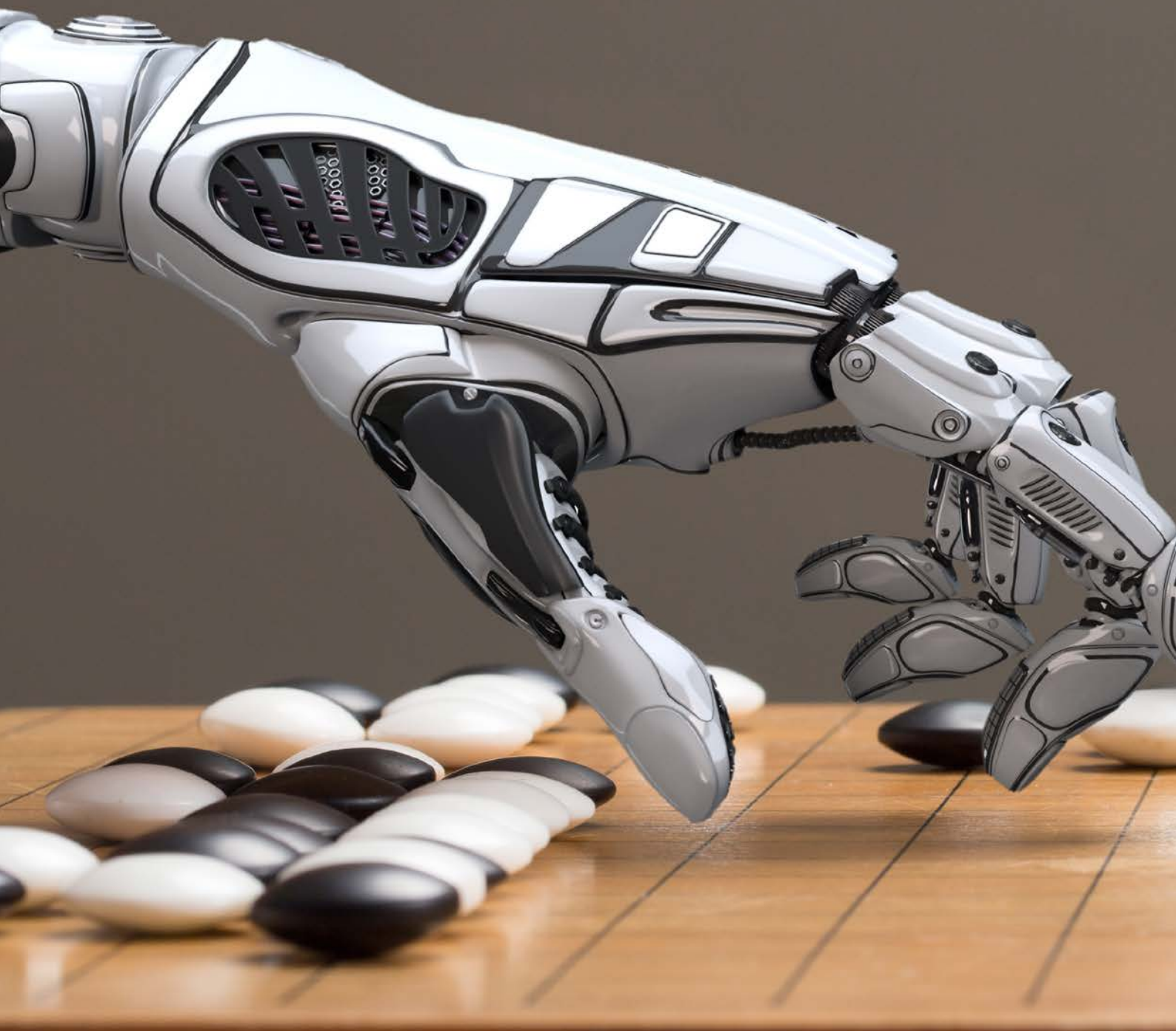
● Un réseau de neurones artificiels à apprentissage profond acquiert sans cesse de l'expertise à partir de nouvelles données.

● Un tel réseau a déjà battu un joueur de go professionnel, d'autres reconnaissent des images ou la parole...

L'AUTEUR



YOSHUA BENGIO
est professeur d'informatique à l'université de Montréal. Il est l'un des pionniers du développement des méthodes d'apprentissage profond.



La révolution de l'APPRENTISSAGE PROFOND

Pour créer une intelligence artificielle, pourquoi ne pas s'inspirer d'une intelligence naturelle ? C'est le pari des réseaux multicouches de neurones artificiels, des « machines » qui apprennent à partir de grandes quantités de données. Leurs succès sont saisissants !

E

n 2016, le programme AlphaGo, mis au point par Google DeepMind, faisait sensation en battant au jeu de go, par le score de 4 à 1, le Sud-Coréen Lee Sedol, réputé le meilleur joueur du monde. Un an après, en octobre 2017, AlphaGo s'est fait humilier avec un score de 100 à 0 par... AlphaGo Zéro, le nouveau logiciel de Google DeepMind!

La différence entre les deux logiciels ? Là où AlphaGo a été nourri par des millions de parties humaines, AlphaGo Zéro s'est contenté des règles du jeu et des positions des pierres sur le plateau. Il a appris en jouant des millions de parties contre lui-même, d'abord au hasard, puis en affinant sa stratégie. En quarante jours, il est devenu le meilleur !

Les deux logiciels ont toutefois un point commun : ils fonctionnent grâce à l'apprentissage profond, aboutissement de décennies de recherche sur l'intelligence artificielle. Aujourd'hui, les machines apprennent grâce à une architecture inspirée de celle du cerveau et d'un nombre considérable de données (AlphaGo Zéro les a générées lui-même). Un retour historique s'impose.

Dans les années 1950, le programme élaboré par Arthur Samuel pour jouer aux dames commence à damer le pion de joueurs de bon niveau. Stupéfaction dans le monde de l'informatique naissante. Le programme en question >



➤ est le premier à inclure un autoapprentissage. L'enthousiasme voire l'euphorie gagne les laboratoires. La conférence de Dartmouth de 1956 consacre la naissance de l'«intelligence artificielle». L'un des pionniers, Marvin Minsky – en 1951, il avait construit la première machine à réseau neuronal, le Snarc – n'hésite pas à déclarer en 1970, dans *Life*: «Dans trois à huit ans nous aurons une machine avec l'intelligence générale d'un être humain ordinaire.» Cet optimisme était un tantinet précipité...

LA CHASSE AU SNARC

Les logiciels de l'époque, par exemple pour aider les médecins dans leur diagnostic, n'ont pas tenu leurs promesses. Les algorithmes étaient trop simples et manquaient des données nécessaires pour parfaire leur apprentissage. La puissance de traitement informatique était également insuffisante pour mener à bien les calculs complexes requis pour imiter, même approximativement, des subtilités de la pensée humaine.

Au milieu des années 2000, le rêve de construire des machines aussi intelligentes que des humains avait presque été abandonné. À cette époque, même le terme d'«intelligence artificielle» semblait avoir déserté le domaine de la science sérieuse. Les chercheurs décrivent la période allant des années 1970 au milieu des années 2000 par l'expression «AI winters», une sorte de long hiver glaciaire de l'intelligence artificielle où tous les espoirs avaient été brisés.

Les choses ont commencé à changer en 2005. Les perspectives dans le domaine de l'intelligence artificielle ont radicalement changé avec l'apprentissage automatique et l'émergence de l'«apprentissage profond», qui revisite le connexionnisme des années 1960 et s'inspire des neurosciences. Les réalisations actuelles de l'apprentissage profond sont à la hauteur des promesses d'antan, et les grandes entreprises des technologies de l'information, y consacrent désormais des milliards de dollars.

Qu'entend-on par apprentissage profond? Il s'agit d'un traitement effectué par un grand nombre de neurones artificiels (imitant de façon très simplifiée les neurones biologiques) qui, par leurs interactions, permettent au système d'apprendre progressivement à partir d'images, de textes ou d'autres données. L'apprentissage repose sur des principes mathématiques généraux. Le résultat de l'apprentissage est une représentation (par exemple, «Cette image contient des éléments différents»), une décision (par exemple «Cette image représente Jeanne Dupont») ou une transformation (par exemple la traduction d'un texte dans une autre langue).

La technique de l'apprentissage profond a transformé la recherche en intelligence artificielle, ranimant des ambitions oubliées pour la vision par ordinateur, la reconnaissance automatique de la parole et la robotique. Les premières applications ont vu le jour en 2012 pour la compréhension de la parole (Siri, qui équipe les iPhone en une illustration). Et peu après sont arrivés des logiciels qui identifient le contenu d'une image, une fonctionnalité qu'intègre maintenant le moteur de recherche de Google.

Tous ceux qui sont frustrés par les menus malcommodes de leur téléphone peuvent apprécier les avantages d'utiliser Siri ou un autre «assistant personnel» sur leur appareil. Et pour ceux qui se rappellent combien la reconnaissance d'objets était mauvaise il y a quelques années seulement, les avancées dans le domaine de la vision artificielle ont été phénoménales: aujourd'hui, moyennant certaines conditions, les ordinateurs savent reconnaître sur des images un chat, une pierre ou des visages presque aussi bien que les humains. Les logiciels d'intelligence artificielle font désormais partie du quotidien de millions d'individus, qui par exemple, n'écrivent plus de messages, mais les dictent à leur téléphone!

Ces progrès ont ouvert la voie à de nouvelles réalisations, à des applications commercialisées, et l'intérêt ne cesse de croître. La concurrence fait rage entre les entreprises pour attirer les jeunes talents... et les moins jeunes. De nombreux professeurs d'université spécialistes du domaine (la majorité, d'après certaines estimations) sont passés du milieu académique à l'industrie, séduits par des équipements de pointe et par de généreux salaires.

Les efforts déployés pour répondre aux défis de l'apprentissage profond ont débouché sur des réussites époustouflantes, les performances d'AlphaGo Zéro en sont l'illustration. Les applications vont s'étendre à d'autres domaines de l'expertise humaine, et pas seulement aux jeux. Par exemple, un algorithme d'apprentissage profond serait capable de diagnostiquer des insuffisances cardiaques sur des images d'IRM aussi bien qu'un cardiologue.

Ce succès récent de l'apprentissage profond a de quoi surprendre quand on se penche sur l'histoire de l'intelligence artificielle. Pourquoi celle-ci a-t-elle buté sur tant d'obstacles dans les décennies précédentes? Parce qu'apprendre des choses nouvelles est, pour une machine, difficile. L'essentiel de la connaissance que nous avons du monde autour de nous n'est pas formalisé dans un langage écrit sous forme d'un ensemble de tâches explicites, ce qui est indispensable pour écrire un programme informatique. C'est pourquoi nous n'avons pas pu directement programmer un ordinateur pour effectuer la plupart des tâches qui nous sont

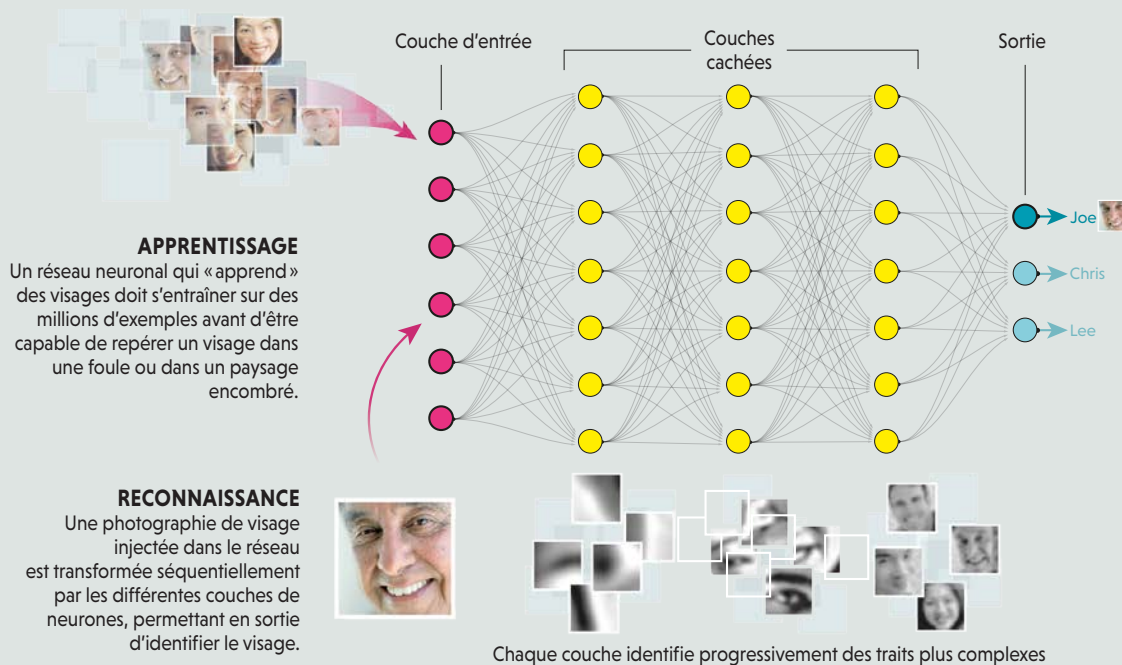
APPRENTISSAGE PROFOND ET ÉTUDE DE L'ADN

Dans le domaine de la génomique, l'accumulation de données nécessite des algorithmes puissants pour en extraire des informations. En 2015, Brendan Frey, de l'université de Toronto, et ses collègues ont développé DeepBind, un logiciel qui recourt à l'apprentissage profond pour analyser la façon dont des protéines se lient à l'ADN et à l'ARN. Les chercheurs détectent ainsi des mutations qui perturbent les processus cellulaires entraînant des maladies.

DES RÉSEAUX DE NEURONES QUI DÉVELOPPENT LEUR EXPERTISE

Les connexions entre neurones dans le cortex cérébral ont inspiré la création d'algorithmes d'apprentissage qui imitent ces liens complexes. On peut apprendre à un tel réseau de neurones artificiels à reconnaître un visage en l'entraînant avec un très grand nombre d'images. Le réseau détermine les traits qui lui permettent de distinguer un visage d'une main, par exemple, et reconnaît la présence de visages dans une image. Il utilise ensuite cette connaissance pour identifier des visages qu'il a déjà vus, même si l'image de la personne

est un peu différente de celle sur laquelle il s'est entraîné. Pour reconnaître un visage dans une image, le réseau commence par analyser les pixels d'une image qui lui est présentée au niveau de la couche d'entrée. Puis il discerne les formes géométriques caractéristiques du visage au niveau de la couche suivante. En remontant la hiérarchie, des yeux, une bouche et d'autres traits du visage apparaissent. Enfin, une forme composite émerge et le réseau tente de « deviner » au niveau de la couche de sortie s'il s'agit du visage de Joe, Chris ou Lee.



évidentes ou faciles, qu'il s'agisse de comprendre la parole, d'interpréter des images ou de conduire une automobile. Les tentatives en ce sens (organiser des ensembles de faits en bases de données élaborées afin de doter les ordinateurs d'un fac-similé d'intelligence) ont rencontré peu de succès.

UNE BONNE DÉCISION

C'est là qu'entre en scène l'apprentissage automatique – le cadre général de l'apprentissage profond. Il se fonde sur des principes généraux, permettant aux systèmes d'utiliser les données disponibles pour apprendre à bien décider, à acquérir de bonnes connaissances et, *in fine*, à rechercher de nouvelles données pour apprendre mieux. Mais qu'est-ce qu'une « bonne » décision ?

Pour les animaux, au regard de l'évolution, une bonne décision optimise les chances de survie et de reproduction. Dans les sociétés humaines, une bonne décision est plus subtile à définir et peut inclure des interactions sociales qui confèrent un statut élevé ou une sensation de bien-être. Pour une voiture autonome la qualité de la prise de décision sera d'autant meilleure que le véhicule reproduira avec fidélité les comportements de bons conducteurs humains.

Les connaissances nécessaires pour prendre une bonne décision dans un contexte particulier ne sont pas nécessairement évidentes et faciles à traduire en langage informatique. Une souris, par exemple, connaît bien son environnement et a un sens inné des endroits où elle doit flirer, sait instinctivement comment bouger ses pattes, trouver de la nourriture, éviter ➤

➤ les prédateurs... Aucun informaticien ne serait capable de spécifier un programme étape par étape pour produire ces comportements. Et pourtant, ces connaissances sont codées dans le cerveau du rongeur.

Avant de créer des ordinateurs capables d'apprendre, les scientifiques doivent répondre à des questions fondamentales, concernant en particulier le partage entre l'inné et l'acquis. Depuis les années 1950, les chercheurs ont étudié et tenté d'affiner les principes généraux qui permettent aux animaux et aux humains (ainsi que, en l'occurrence, aux machines) d'acquérir des connaissances par expérience. L'apprentissage automatique vise à établir des procédures, des algorithmes d'apprentissage, qui confèrent à une machine la capacité d'apprendre à partir d'exemples qu'on lui présente.

La science de l'apprentissage automatique fait face à un résultat négatif formel surnommé *No free lunch* (« Rien n'est gratuit ») : si tous les problèmes sont équiprobables, il n'existe pas d'algorithme universel, c'est-à-dire meilleur que tous les autres sur l'ensemble des problèmes. Il faut donc élaborer des algorithmes distincts pour s'attaquer à différentes catégories de problèmes (par exemple reconnaître un coucher de soleil ou traduire un texte en ourdou). L'apprentissage automatique repose ainsi sur deux piliers, théorique et expérimental : un problème réel ne satisfait pas toujours les hypothèses théoriques de l'algorithme, et la validation de l'algorithme sur des données réelles est toujours nécessaire.

Or il semble que notre cerveau incorpore des algorithmes généraux qui nous permettent d'apprendre une multitude de tâches auxquelles l'évolution n'avait pas préparé nos ancêtres : jouer aux échecs, construire des ponts ou faire de la recherche en intelligence artificielle.

Ces compétences supplémentaires suggèrent que l'intelligence humaine pourrait encore être source d'inspiration pour créer des machines dotées d'une forme d'intelligence générale. C'est exactement pourquoi les développeurs ont adopté le modèle du cerveau comme guide pour concevoir des systèmes intelligents sous la forme de réseaux de neurones.

DES CERVEAUX VIRTUELS

La principale unité de calcul du cerveau est le neurone. Cette cellule transmet un signal sous la forme d'une impulsion électrique qui se propage jusqu'à la synapse, la zone de communication avec un autre neurone. Des molécules nommées neurotransmetteurs sont libérées, puis réabsorbées par le neurone cible, qui prend alors le relais. La propension d'un neurone à transmettre un signal vers un autre neurone est la force synaptique. À mesure qu'un neurone « apprend », sa force synaptique augmente, et il a plus de chances

d'envoyer des signaux à ses voisins quand il est stimulé par une impulsion électrique.

Les neurosciences ont influencé l'émergence des réseaux de neurones artificiels, où ces éléments sont connectés de façon matérielle ou logicielle. Les premiers programmeurs de cette sous-discipline de l'intelligence artificielle, connue sous le nom de connexionnisme, postulaient que les réseaux neuronaux seraient capables d'apprendre des tâches complexes en modifiant progressivement les connexions entre neurones, de telle façon que les schémas d'activité neuronale coderaient le contenu des données livrées sous forme d'images, de sons... À chaque exemple soumis au réseau, le processus d'apprentissage se poursuivrait en modifiant les forces synaptiques entre neurones connectés. Ces valeurs convergeraient petit à petit vers la configuration du réseau qui représente le mieux, par exemple, les images de couchers de soleil.

Les réseaux de neurones actuels sont des versions améliorées des travaux pionniers issus du connexionnisme. Cependant, les algorithmes d'apprentissage correspondants requièrent une participation active de l'homme. La plupart d'entre eux utilisent un apprentissage supervisé, où chaque exemple d'apprentissage est accompagné d'une étiquette indiquant le sujet de l'apprentissage, tel que « coucher de soleil ».

Dans ce cas, l'objectif de l'algorithme d'apprentissage supervisé est, à partir d'une donnée d'entrée constituée par une photographie, de produire comme sortie le nom de l'objet central de l'image. L'opération mathématique qui transforme une entrée en une sortie est une fonction. Les valeurs numériques qui définissent cette fonction, telles que les forces synaptiques, constituent une solution de la tâche d'apprentissage.

**Dans trois à huit ans
nous aurons
une machine avec
l'intelligence
générale d'un être
humain ordinaire...
pouvait-on lire
en 1970**

Il serait facile pour un système d'apprendre par cœur les réponses correctes sur les exemples connus (il suffit d'une mémoire suffisante). Mais cela n'a pas grand intérêt. Même si on accumule des millions d'images de coucher de soleil, le nombre d'images possibles d'une telle scène est infini. L'objectif de l'apprentissage est d'être capable de donner de bonnes réponses sur de nouveaux exemples, inconnus du système, en généralisant les exemples connus. Le bon niveau de généralisation dépend du contexte. Ainsi, on peut vouloir reconnaître la notion d'arbre en général, mais on peut aussi vouloir distinguer les chênes des hêtres, ou vouloir reconnaître le hêtre d'un jardin donné...

Un tel algorithme doit aussi s'appuyer sur certaines hypothèses relatives aux données et à ce que pourrait être une solution possible à un problème donné. Ainsi, le logiciel doit intégrer comme principe que si des données d'entrée d'une fonction particulière sont semblables, leur sortie ne devrait pas être radicalement différente: modifier quelques pixels sur une image de chat ne devrait pas transformer l'animal en chien.

De telles hypothèses faites sur des images sont utilisées par les réseaux de neurones dits convolutifs. Ces programmes sont à l'origine du renouveau de l'intelligence artificielle. Les réseaux neuronaux convolutifs utilisés dans l'apprentissage profond comprennent de nombreuses couches de neurones organisées de sorte que le logiciel sera robuste vis-à-vis des changements dans l'objet qu'il tente d'analyser. Il reconnaîtra l'élément même s'il a un peu bougé. Ainsi, un réseau bien entraîné identifiera un visage quel que soit l'angle de vue.

UNE STRUCTURE MULTICOUCHE

La configuration d'un réseau convolutif s'inspire de la structure en couches multiples du cortex visuel, c'est-à-dire la partie de notre cerveau qui reçoit les signaux des yeux. Ces nombreuses couches de neurones virtuels justifient le qualificatif de «profond» donné à ces réseaux et confèrent à ces derniers une meilleure capacité à appréhender le monde environnant.

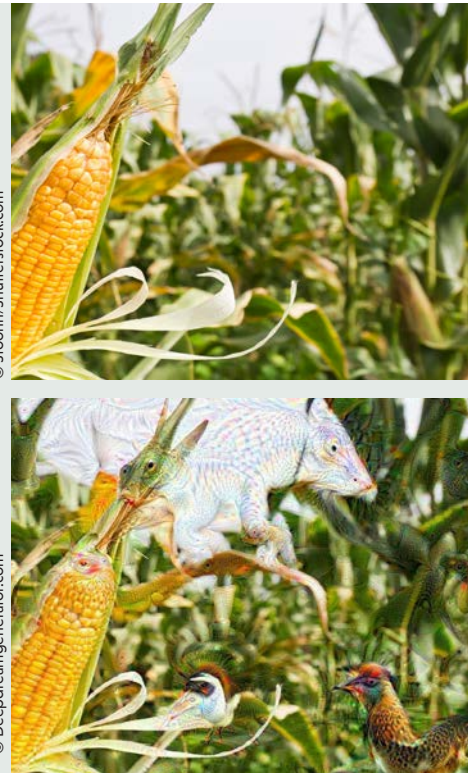
L'apprentissage profond est devenu envisageable il y a une dizaine d'années grâce à certaines innovations, alors que l'intérêt pour l'intelligence artificielle était au plus bas. Un organisme canadien financé par le gouvernement et par des fonds privés, l'Icra (Institut canadien de recherches avancées), a contribué à ranimer la flamme en soutenant un programme dirigé par Geoffrey Hinton, de l'université de Toronto, et dont faisaient partie Yann LeCun, de l'université de New York et du centre de recherche de Facebook à Paris, Andrew Ng, de l'université de Stanford, Bruno Olshausen, de l'université de Californie à Berkeley, moi-même et

APPRENTISSAGE PROFOND ET ART PSYCHÉDÉLIQUE

L'apprentissage profond est performant pour la reconnaissance d'objets dans une image, mais il est en réalité assez difficile de comprendre ce qui se passe au cours de l'analyse à chaque niveau du réseau de neurones. La société Google a développé le logiciel DeepDream qui utilise l'apprentissage profond. Les chercheurs peuvent analyser le résultat à différentes étapes du processus en demandant au logiciel de faire ressortir ce qu'il «voit» dans l'image. Ainsi, telles des paréidolies – l'illusion de voir, par exemple, un chien dans un nuage qui en a vaguement la forme –, le logiciel, qui s'est, par exemple, entraîné à reconnaître des images d'animaux, verra des oiseaux et des rongeurs dans une photographie d'un champ de maïs. Chacun peut créer ses propres images psychédéliques sur le site Deepdreamgenerator.com.

© Sironn/Shutterstock.com

© Deepdreamgenerator.com



plusieurs autres. En raison du scepticisme ambiant, il était difficile à cette époque de publier des articles sur le sujet et même de persuader des étudiants de faire leur thèse dans ce domaine. Mais nous étions convaincus qu'il fallait persévérer dans cette voie.

La réticence vis-à-vis de l'intelligence artificielle découlait initialement de l'idée que la création de réseaux neuronaux était sans espoir à cause de la difficulté à optimiser leur performance pour apprendre efficacement.

L'optimisation est une branche des mathématiques qui vise à trouver la meilleure combinaison de paramètres, ici les poids synaptiques des neurones virtuels, pour atteindre un certain objectif. Lorsque la relation qui lie les paramètres et l'objectif est assez simple (on dit que l'objectif est une «fonction convexe» des paramètres) alors on peut améliorer progressivement les paramètres et l'on parle d'optimisation convexe. La procédure d'entraînement est répétée jusqu'à ce que les paramètres s'approchent aussi près que possible des valeurs qui produisent le meilleur résultat. On cherche en fait un minimum global, à savoir le jeu de paramètres qui permet d'atteindre la valeur la plus basse (et la meilleure) de l'écart à l'optimum.

En général, la situation n'est pas aussi simple, la fonction reliant les paramètres à l'objectif qu'ils atteignent n'étant pas convexe. Or l'optimisation non convexe représente un défi bien plus difficile. De nombreux chercheurs >

➤pensaient qu'il n'était pas possible de le relever. L'algorithme d'apprentissage pourrait en effet se retrouver coincé dans un minimum local, d'où il ne pourrait sortir en ajustant les paramètres par petites touches.

Or avec mes collègues, nous avons montré que si le réseau de neurones est assez grand, le problème des minima locaux est fortement réduit. Dans un tel réseau, la plupart des minima locaux correspondent à un niveau d'apprentissage quasi équivalent à celui du minimum global.

Bien que les problèmes théoriques de l'optimisation soient, en principe, résolubles, la construction de grands réseaux comportant plus de deux ou trois couches avait souvent échoué. À partir de 2005, les recherches financées par l'Icra ont commencé à porter leurs fruits et les obstacles ont été surmontés progressivement. En 2006, nous sommes parvenus à entraîner des réseaux neuronaux plus profonds en utilisant une technique qui opérait couche par couche.

Par la suite, en 2011, nous avons trouvé un moyen d'entraîner des réseaux encore plus profonds (avec davantage de couches de neurones virtuels) en modifiant les opérations effectuées par chacun des neurones, ce qui les rapprochait davantage des neurones biologiques dans leur façon de traiter l'information. Nous avons aussi découvert qu'en injectant du bruit aléatoire dans les signaux transmis entre les neurones au cours de l'apprentissage – comme ce qui a cours dans le cerveau –, les réseaux apprenaient mieux à identifier correctement une image ou un son.

Par ailleurs, deux facteurs essentiels ont contribué au succès des techniques d'apprentissage profond. Le premier est l'augmentation d'un facteur 10 de la puissance de calcul des ordinateurs grâce aux processeurs dédiés au traitement d'image, ce qui a permis d'entraîner des réseaux plus grands en un temps raisonnable. Le second facteur est que l'on avait désormais accès à d'énormes bases de données étiquetées, avec lesquelles les algorithmes d'apprentissage ont pu s'exercer à reconnaître un « chat », par exemple.

Une autre raison aux récents succès de l'apprentissage profond est que le réseau est capable d'apprendre à effectuer un ensemble de calculs qui construisent ou analysent pas à pas une image, un son ou un autre type de donnée. Avec une profondeur suffisante, ces dispositifs excellent dans beaucoup de tâches de reconnaissance visuelle ou auditive. Des travaux théoriques et expérimentaux récents ont d'ailleurs montré qu'il est impossible d'effectuer efficacement certaines de ces opérations mathématiques sans un réseau assez profond.

Mais que se passe-t-il au cœur du réseau de neurones profond? Chaque couche transforme son entrée et produit une sortie qui est envoyée à la couche suivante (voir l'encadré

page 45). Les premières couches se concentrent sur les détails aux échelles les plus petites, puis les couches suivantes agrandissent les échelles considérées. Plus les couches sont profondes, plus elles vont représenter des concepts abstraits.

RECONNAÎTRE DES VIDÉOS

Jusqu'à récemment, les réseaux de neurones artificiels se distinguaient en grande partie par leur capacité à effectuer des tâches telles que la reconnaissance de motifs dans des images statiques. Mais d'autres types de réseaux neuronaux affichent des résultats intéressants sur des événements dynamiques. Les « réseaux neuronaux récurrents », par exemple, ont démontré leur capacité à effectuer une séquence de calculs pour traiter la parole ou une vidéo. Les données séquentielles sont constituées d'unités (phonèmes ou mots dans le cas de la parole) qui se suivent. La façon dont les réseaux neuronaux récurrents traitent leurs entrées ressemble un peu au fonctionnement du cerveau. Les signaux qui se propagent parmi les neurones changent constamment à mesure que les entrées sensorielles sont traitées.

Les réseaux récurrents parviennent à prédire quel sera le mot suivant dans une phrase et peuvent ainsi produire une nouvelle séquence de mots, l'un après l'autre. Ils peuvent aussi s'atteler à des tâches plus complexes. Après avoir « lu » tous les mots d'une phrase, le réseau est à même de deviner le sens de la phrase entière. Un réseau récurrent distinct peut alors utiliser le traitement sémantique du premier réseau pour traduire la phrase dans une autre langue.

La recherche sur les réseaux neuronaux récurrents a connu sa propre période de stagnation à la fin des années 1990 et au début des années 2000. Mes propres travaux théoriques suggéraient qu'ils rencontreraient des difficultés à apprendre à récupérer de l'information du passé lointain, à savoir les premiers éléments de la séquence en cours de traitement. Mais certains de ces problèmes ont été résolus grâce à des techniques consistant à stocker l'information pour qu'elle persiste durablement. Ces réseaux neuronaux utilisent la mémoire temporaire (mémoire cache) d'un ordinateur pour traiter des informations multiples et dispersées, comme des idées contenues dans différentes phrases réparties au fil d'un document.

Le retour en force des réseaux neuronaux à apprentissage profond après le long sommeil de l'intelligence artificielle n'est pas uniquement un triomphe technique. C'est aussi une leçon de sociologie des sciences. En particulier, il souligne la nécessité de soutenir des idées qui font fi du *statu quo* technologique et d'encourager la tenue d'un portefeuille de recherches diversifié, laissant de la place à des champs provisoirement passés de mode. ■

BIBLIOGRAPHIE

D. SILVER ET AL., *Mastering the game of Go without human knowledge*, *Nature*, vol. 550, pp. 354-359, 2017.

Y. LECUN ET AL., *Deep learning*, *Nature*, vol. 521, pp. 436-444, 2015.

Y. BENGIO ET AL., *Representation learning: A review and new perspectives*, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 35(8), pp. 1798-1828, 2013.

ImageNet classification with deep convolutional neural networks, 26th Annual Conference on Neural Information Processing Systems (NIPS 2012), 2012.

SUR LE WEB

Cours de Yann LeCun au Collège de France, 2015-2016: <https://www.college-de-france.fr/site/yann-lecun/index.htm>



« Quand les “data scientists” nous simplifient la vie »

En ce début de XXI^e siècle, la science des données s’imisce tous les jours un peu plus dans notre quotidien. Reconnaissance vocale, moteurs de recherche, publicités ciblées, recommandations d’achats... Elle est bien sûr déjà omniprésente dans nos activités numériques. Mais cette *data science* gagne aussi progressivement tous les secteurs d’activités économiques. Décoller en A380, emprunter sereinement son train quotidien ou bien encore se remettre d’une catastrophe naturelle : voici trois cas très concrets où des algorithmes aussi discrets qu’efficaces sont déjà à la manœuvre. Autant d’applications qui nécessitent les compétences de *data scientists* chevronnés... Un secteur d’activité florissant qui pourrait créer trois millions d’emplois en Europe ! Et grâce au nombre et à la qualité de ses formations, la France est bien placée pour ne pas laisser filer sa part du gâteau... Rien que sur l’Hexagone, le recrutement de 130 000 *data scientists* sera nécessaire d’ici à 2020. En ferez-vous partie ?

SOMMAIRE

- Décoller à l’heure p. II
- Des sinistrés pris en charge plus vite p. III
- Améliorer l’information temps réel p. IV
- Et si vous deveniez *data scientist* ? p. V

Cahier spécial réalisé en partenariat avec

