

véhicules motorisés, l'ordinateur doit analyser des images variées et changeantes : où est le bord de cette rue jonchée de feuilles d'arbres ? Quelle est la nature de cette zone noire à 50 mètres au centre de la chaussée, un trou dangereux ou seulement une tache d'huile inoffensive ? Etc.

Conduire une voiture avec nos méthodes nécessiterait la mise au point de techniques d'analyse d'images bien plus subtiles que celles que nous savons programmer aujourd'hui. Aussi, les systèmes de pilotage automatisé tels que ceux de la firme *Google* « conduisent » tout à fait différemment des humains. Ces *Google cars* exploitent en continu un système GPS de géolocalisation très précis et des « cartes » indiquant de manière bien plus détaillée que toutes les cartes habituelles, y compris celles de *Google maps*, la forme et le dessin des chaussées, la signalisation routière et tous les éléments importants de l'environnement. Les voitures *Google* exploitent aussi des radars embarqués, des lidars (*light detection and ranging*, systèmes optiques créant une image numérique en trois dimensions de l'espace autour de la voiture) et des capteurs sur les roues.

Ayant déjà parcouru plusieurs centaines de milliers de kilomètres sans accident, ces voitures sont un succès de l'intelligence artificielle, et ce même si elles sont incapables de réagir à des signes ou injonctions d'un policier au centre d'un carrefour, et qu'elles s'arrêtent parfois brusquement lorsque des travaux sont en cours sur leur chemin. Par prudence sans doute, les modèles destinés au public présentés en mai 2014 roulent à 40 kilomètres par heure au plus.

Avec ces machines, on est loin de la méthode de conduite d'un être humain. Grâce à sa capacité à extraire de l'information des images et son intelligence générale, le conducteur humain sait piloter sur un trajet jamais emprunté, sans carte, sans radar, sans lidar, sans capteur sur les roues et il n'est pas paralysé par un obstacle inopiné !

Les questions évoquées jusqu'ici n'exigent pas la compréhension du langage écrit ou parlé. Pourtant, contrairement aux annonces de ceux qui considéraient le langage comme une source de difficultés insurmontables pour les machines, des succès remarquables

ont été obtenus dans des tâches exigeant une bonne maîtrise des langues naturelles.

L'utilisation des robots-journalistes inquiète car elle est devenue courante dans certaines rédactions telles que celles du *Los Angeles Times*, de *Forbes* ou de *Associated Press*. Pour l'instant, ces automates-journalistes se limitent à convertir des résultats (sportifs, ou économiques par exemple) en courts articles.

Des robots journalistes

Il n'empêche que, parfois, on leur doit d'utiles traitements. Ainsi, le 17 mars 2014, un tremblement de terre de magnitude 4,7 se produisit à 6 h 25 en heure locale au large de la Californie. Trois minutes après la secousse, un petit article d'une vingtaine de lignes était automatiquement publié sur le site Internet du *Los Angeles Times* donnant des informations sur l'événement : lieu de l'épicentre, magnitude, heure, comparaison avec d'autres secousses récentes. L'article exploitait des données brutes fournies par le *US-Geological Survey Earthquake Notification Service* et résultait d'un algorithme dû à Ken Schwencke, un journaliste programmeur. D'après lui, ces méthodes ne conduiraient à la suppression d'aucun emploi, mais rendraient au contraire le travail des journalistes plus intéressant. Il est vrai que ces programmes sont pour l'instant confinés à la rédaction d'articles brefs exploitant des données factuelles faciles à traduire en petits textes, qu'un humain ne rédigerait sans doute pas mieux.

Un autre exemple inattendu de rédaction automatique d'articles concerne les encyclopédies *Wikipedia* en suédois et en filipino, l'une des deux langues officielles aux Philippines (l'autre est l'anglais). Le programme *Lsjbot* mis au point par Sverker Johansson a en effet créé plus de deux millions d'articles de l'encyclopédie collaborative et est capable d'en produire 10 000 par jour. Ces pages engendrées automatiquement concernent des animaux ou des villes et proviennent de la traduction, dans le format imposé par *Wikipedia*, d'informations disponibles dans des bases de données déjà informatisées. L'exploit a été salué, mais aussi critiqué. Pour se justifier, S. Johansson indique que ces

pages peu créatives sont utiles et fait remarquer que le choix des articles de l'encyclopédie *Wikipedia* est biaisé : il reflète essentiellement les intérêts des jeunes blancs, de sexe masculin et amateurs de technologies. Ainsi, l'encyclopédie *Wikipedia* suédoise comporte 150 articles sur les personnages du livre de Tolkien, *Le Seigneur des Anneaux*, et seulement une dizaine sur des personnes réelles liées à la guerre du Vietnam : « Est-ce vraiment le bon équilibre ? », demande-t-il.

S. Johansson projette d'engendrer une page par espèce animale recensée, ce qui ne semble pas stupide. Pour lui, ces méthodes doivent être généralisées, mais il pense que *Wikipedia* a besoin aussi de rédacteurs qui écrivent de façon plus littéraire que *Lsjbot* et soient capables d'exprimer des sentiments, « ce que ce programme ne sera jamais capable de faire ».

Beaucoup plus complexe, et méritant mieux l'utilisation de l'expression « intelligence artificielle », est le succès du programme *Watson* d'IBM au jeu télévisé *Jeopardy* ! (voir l'encadré 1).

La mise au point de programmes résolvant les mots-croisés aussi bien que les meilleurs humains confirme que l'intelligence artificielle réussit à développer des systèmes aux étonnantes performances linguistiques et oblige à reconnaître qu'il faut cesser de considérer que le langage est réservé aux humains. Malgré ces succès, on est loin de la perfection : pour s'en rendre compte, il suffit de se livrer à un petit jeu avec le système de traduction automatique en ligne de *Google*. Une phrase en français, [a], est traduite en anglais, [b], puis retraduite en français, [c]. Parfois cela fonctionne parfaitement, on retrouve la phrase initiale ou une phrase équivalente. Parfois, le résultat est catastrophique, comme ici :

[a] La subtilité manque aux machines qui semblent stupides.

[b] The lack subtlety machines that seem stupid.

[c] Les machines manque de subtilité qui semblent stupides.

Nous sommes très loin aujourd'hui de la mise au point de dispositifs informatiques susceptibles de passer le « test de Turing » conçu en 1950. Alan Turing voulait éviter

de discuter de la nature de l'intelligence et, plutôt que d'en rechercher une définition, proposait de considérer qu'on aura réussi à mettre au point des machines intelligentes lorsque leur conversation sera indiscernable de celle des humains.

Passer le test de Turing

Pour tester cette indiscernabilité, il suggérait de faire dialoguer par écrit avec la machine une série de juges qui ne sauraient pas s'ils mènent leurs échanges avec une machine tentant de se faire passer pour un humain ou avec un humain véritable. Lorsque les juges ne pourront plus faire mieux que répondre au hasard pour indiquer qu'ils ont eu affaire à un homme ou une machine, le test sera passé. Concrètement, faire passer

le test de Turing à un système informatique S consiste à réunir un grand nombre de juges, à les faire dialoguer aussi longtemps qu'ils le souhaitent avec des interlocuteurs choisis pour être une fois sur deux un humain, et une fois sur deux le système S; les experts indiquent, quand ils le souhaitent, s'ils pensent avoir échangé avec un humain ou une machine. Si l'ensemble des experts ne fait pas mieux que le hasard, donc se trompe dans 50 % des cas ou plus, alors le système S a passé le test de Turing.

Turing, optimiste, pronostiqua qu'on obtiendrait une réussite partielle au test en l'an 2000, les experts dialoguant 5 minutes et prenant la machine pour un humain dans 30 % des cas au moins. Turing avait, en gros, vu juste : depuis quelques années, la version partielle du test a été passée, sans qu'on

puisse prévoir quand sera passé le test complet, sur lequel Turing restait muet.

Le test partiel a par exemple été passé le 6 septembre 2011 à Guwahati, en Inde, par le programme *Cleverbot* créé par l'informaticien britannique Rollo Carpenter. Trente juges dialoguèrent pendant quatre minutes avec un interlocuteur inconnu qui était dans la moitié des cas un humain et dans l'autre moitié des cas le programme *Cleverbot*. Les juges et les membres de l'assistance (1334 votes) ont considéré le programme comme humain dans 59,3 % des cas. Notons que les humains ne furent considérés comme tels que par 63,3 % des votes.

Plus récemment, le 9 juin 2014 à la *Royal Society* de Londres, un test organisé à l'occasion du soixantième anniversaire de la mort de Turing permit à un programme

2. Le concours de compression de Marcus Hutter

Face à une série de données, l'intelligence permet de prévoir la suite mieux que ne le ferait le hasard. On montre que cette capacité de prévision est équivalente à la capacité de compresser les données reçues. À condition d'imaginer des jeux de données très variés, cette conception de l'intelligence est considérée par certains comme absolue.

C'est elle qui sert d'ailleurs de base à la « théorie générale de l'intelligence » de Marcus Hutter. Ce dernier, pour illustrer l'idée et parce qu'il est convaincu que la compression de données est nécessaire et suffisante pour caractériser et mesurer l'intelligence d'un être animal, humain ou mécanique, a créé un concours fondé sur une épreuve de compression de données. Ce concours est ouvert à tous et permet de gagner tout ou partie d'une somme de 50 000 euros que M. Hutter a lui-même engagée.

Le prix (nommé *Hutter prize*, voir <http://prize.hutter1.net/>) propose une épreuve unique de compression d'un texte numé-

rique énorme. Un fichier de 100 millions de caractères extrait de l'encyclopédie *Wikipedia* est proposé aux candidats et doit être compressé au mieux par les programmes qui sont soumis. Au départ, le fichier, grâce à une méthode de compression classique, a été réduit à 18 324 887 caractères. Cela correspond à un gain d'environ 81 %. Chaque candidat proposant un programme réalisant une amélioration de N % par rapport au gagnant précédent remporte la somme de $(N/100) \times 50\,000$ euros. Si par exemple vous concevez et programmez un compresseur qui fait gagner 5 % par rapport au dernier gagnant, vous

emportez 2500 euros. La compression dont il s'agit ici est bien sûr une compression sans perte : à partir du fichier compressé, le programme de décompression (associé au programme de compression) doit reconstituer exactement les 100 millions de caractères du fichier initial proposé par M. Hutter. Des limites concernant la taille du programme de compression et son temps de calcul sont imposées (voir les détails sur le site Internet décrivant le prix).

Depuis le lancement du concours, plusieurs programmes ont réussi à améliorer le taux de compression. On en est aujourd'hui à un fichier compressé de 15 949 688 caractères. Claude Shannon, le créateur de la théorie de l'information, a évalué que le langage naturel porte en gros un bit d'information par caractère, ce qui correspond dans notre cas à un fichier



UNE COMPRESSION RÉALISÉE
par le sculpteur César.

compressé de 12 500 000 caractères. Il y a donc encore de la marge, d'autant que l'exploitation fine des régularités de notre monde reflétée dans l'encyclopédie *Wikipedia* permettrait peut-être une compression encore meilleure que celle évaluée par Shannon.