



LA CHRONIQUE DE
G RALD BRONNER

professeur de sociologie   l'universit  Paris-Diderot

VLADIMIR POUTINE ET BARBRA STREISAND

**  l' re d'Internet et des r seaux sociaux, l'interdiction
d'une image risque de lui faire beaucoup de publicit .
Plusieurs c l brit s l'ont appris   leurs d pens...**



Cette image
a  t  interdite par
la justice russe.
Avec un r sultat
inverse de celui
recherch  :
l'image s'est
diffus e
massivement.

On ne songerait pas forc ment   associer ces deux personnages. L'un est une star mondiale, actrice, chanteuse, r alisatrice bien connue. L'autre est une importante figure politique de notre temps. Leader sulfureux de la Russie, il est souvent accus  d'exercer le pouvoir de mani re autoritaire. Et comment interpr ter autrement la fa on dont son minist re de la Justice a fait interdire une liste d'images ou de productions consid r es comme extr mistes (4074   ce jour) ? Parmi celles-ci figure une repr sentation du pr sident de la Russie f minis  par un maquillage (*ci-dessus*), image jug e infamante selon une d cision d'un tribunal rendue en mars 2017. Celui-ci a consid r  que l'image en question sugg re une « orientation sexuelle non conforme du pr sident russe ».

C'est l  que le destin de l'actrice am ricaine et du leader russe se rencontrent. La premi re a donn  son nom   un effet bien connu en communication : l'effet Streisand. De quoi s'agit-il ? En 2003, un conflit opposa la star hollywoodienne  

Kenneth Adelman, qui avait pris en photo la c te de Malibu afin, pr tendit-il, d' tudier l' rosion du littoral. Le probl me  tait que la somptueuse villa de Barbara Streisand, situ e sur cette c te,  tait clairement visible sur le clich . Arguant de la loi antipaparazzi de Californie, elle d cida d'attaquer l'ind licat en justice en esp rant r duire la diffusion du clich .



Villa de Barbra Streisand : un incident prototypique d'un ph nom ne de communication



Mal lui en prit : c'est tout l'inverse qui se produisit. L'affaire s' bruitant peu   peu, plusieurs sites internet reprirent la photographie, qui fut vue 420000 fois dans le mois suivant.

Cet incident est d sormais consid r  comme prototypique d'un ph nom ne de communication ressemblant   la parabole du pompier incendiaire et que l'on nomme donc l'effet Streisand. Il arrive que les efforts faits

pour emp cher la diffusion d'une information y contribuent, au contraire. Et c'est pr cis ment ce qui est en train d'arriver   Vladimir Poutine avec le photomontage   la Andy Warhol. Son auteur, un certain Aleksandr Tsvetkov, vient d' tre envoy  en centre de soins psychiatriques, mais, pour le pouvoir russe, le mal est fait : l'image se r pand partout et devient m me une bann re pour les opposants et les critiques du dirigeant russe.

Ce dernier n'est pas le seul    tre victime de l'effet Streisand. Pour prendre un exemple parmi d'autres, mais venant des  tats-Unis cette fois, l'artiste contemporaine Ilma Gore a vu sa toile interdite d'exposition parce qu'elle repr sentait Donald Trump dans le plus simple appareil, et de fa on peu flatteuse. Cette m saventure a assur  une belle publicit    son  uvre, qui a  t  mise en vente dans une galerie londonienne pour... un million de livres !

De fa on g n rale, les effets Streisand ont tendance   se multiplier avec la d r gulation du march  de l'information qui caract rise un monde libre et connect . D'une part, parce qu'un tel march  rend plus difficile le contr le des producteurs d'informations. D'autre part, parce qu'une information interdite prend, au moins symboliquement, une valeur de raret  sur un march  devenu ultraconcurrentiel. Une interdiction assure ainsi une plus-value   tous les acteurs cherchant d'une fa on ou d'une autre   capter un peu de temps de cerveau disponible.

Le rire et la moquerie ont toujours constitu  une forme de danger pour toute autorit , car ils d voient la vanit  de ses expressions symboliques. Longtemps – depuis Attila au moins –, ceux qui incarnent le pouvoir se sont fait accompagner de bouffons, sans doute en partie pour contr ler les expressions du rire. On dit m me que le dernier des fous du roi, l'Ang ly, au service de Louis XIII et de Louis XIV,  tait si craint qu'on le payait pour  viter d' tre l'objet de ses terribles railleries. Il n'y aurait pas aujourd'hui de bourses assez fournies pour acheter le silence des millions de r ieurs qui attendent sur la Toile. ■

Au-delà du TEST DE TURING

Les machines peuvent-elles penser? C'est en posant cette question que débute l'article «Computing machinery and intelligence», publié en 1950 par le mathématicien anglais Alan Turing, casseur du code Enigma utilisé par les Allemands durant la Seconde Guerre mondiale et pionnier de la science informatique. Au lieu de tenter de répondre à la question initiale qui soulève des difficultés conceptuelles (en particulier, que signifie «penser»?), Turing proposait un «jeu d'imitation» qui, s'il était gagné par une machine, prouverait que celle-ci a des capacités intellectuelles proches de celles d'un humain.

En deux mots, une machine gagne le jeu d'imitation si elle parvient, à travers un échange de textes avec un juge humain, à convaincre ce dernier qu'il a affaire à un interlocuteur en chair et en os. Ce «test de Turing» est devenu dans les décennies suivantes une sorte de guide pour les domaines de l'intelligence artificielle et de la robotique – un défi qui, s'il était relevé, signifierait que la machine a atteint le même niveau d'intelligence que les hommes.

À ce jour, aucun programme informatique ni robot n'a vraiment passé avec succès le test de Turing. Mais comme l'expliquent les chercheurs Gary Marcus et Laurence Devillers dans les pages qui suivent, les progrès scientifiques et techniques ont fait apparaître que l'on ne

peut plus considérer le test de Turing comme l'objectif ultime à atteindre en matière d'intelligence artificielle.

Il y a à cela plusieurs raisons. D'une part, le jeu d'imitation de Turing ne teste que la capacité d'un système artificiel à bernier un examinateur. Ce test, souligne Gary Marcus dans son article («*Machine, es-tu intelligente?*», pages 26-31), doit être complété par des épreuves nouvelles portant sur diverses capacités intellectuelles, et des propositions concrètes existent déjà.

D'autre part, la notion même d'intelligence recouvre aujourd'hui un vaste éventail de capacités qui ne sont pas seulement d'ordre intellectuel, mais aussi d'ordre émotionnel. Or avec l'apparition de systèmes informatiques ou de robots capables de dialoguer oralement avec nous, de détecter notre état émotif, d'exprimer de l'empathie, la question de l'évaluation de ces aptitudes émotionnelles se pose avec acuité, explique Laurence Devillers («*Tester les robots pour mieux vivre avec*», pages 32-38). Dans un monde où nous côtoierons quotidiennement des robots évolutifs, la question de tous ces tests est aussi un enjeu de société: il en va de la sécurité et de la santé mentale des humains. ■

MAURICE MASHAAL

L'ESSENTIEL

> Une machine soumise au test de Turing tente, par ses réponses, de convaincre un interrogateur qu'elle est un humain.

> Ce test a longtemps été considéré comme la meilleure façon de jauger une intelligence artificielle.

> Le test de Turing a cependant mal vieilli. Le réussir est davantage affaire de tromperie que de véritable intelligence. Pour les experts, le test de Turing est à remplacer par une batterie d'épreuves qui évalueraient l'intelligence de la machine sous différents angles.

> Une machine réellement intelligente devrait pouvoir comprendre des affirmations ambiguës, assembler un meuble vendu en pièces détachées, réussir un examen scolaire... Étant donné la difficulté de ces tâches, l'homme n'est pas près d'être détrôné sur le plan de l'intellect.

L'AUTEUR



GARY MARCUS
directeur d'Uber AI Labs
et professeur de psychologie et
de neurosciences à l'université
de New-York, aux États-Unis

Machine, es-tu intelligente?

ON A LONGTEMPS CHERCHÉ LA RÉPONSE À CETTE QUESTION À L'AIDE DU FAMEUX TEST DE TURING. MAIS CE TEST, IMAGINÉ EN 1950, A MONTRÉ SES LIMITES. PAR QUOI LE REMPLACER ? PAR DES BATTERIES D'ÉPREUVES DE DIFFÉRENTES NATURES, DISENT LES SPÉCIALISTES.

En 1950, le mathématicien britannique Alan Turing imagine une expérience de pensée devenue célèbre et souvent considérée comme le meilleur moyen de cerner l'intelligence d'une machine. Il l'appela le « jeu de l'imitation », mais on la connaît généralement sous le nom de « test de Turing ».

Anticipant les *chat bots* d'aujourd'hui, ou agents conversationnels, des programmes informatiques qui se font passer pour des humains, Turing envisageait un échange textuel au cours duquel une machine tenterait de tromper un examinateur et de l'amener à penser qu'elle est un humain, en répondant à des questions relatives à de la poésie et en commettant délibérément des erreurs de calcul.

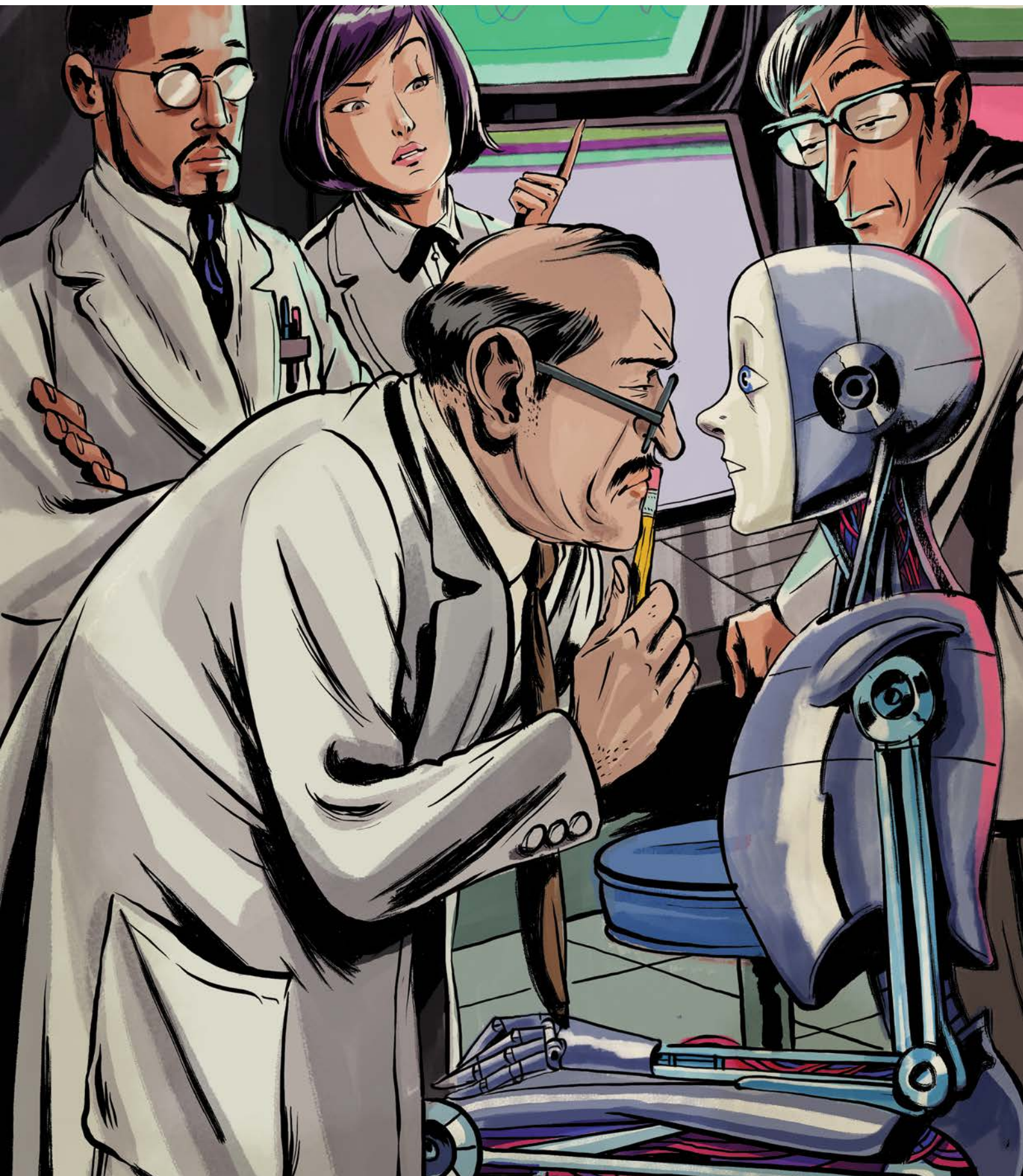
De nos jours, le grand public voit souvent le test de Turing comme une sorte de Rubicon, dont le franchissement signifierait que la machine est devenue l'égale de l'homme, du moins sur le plan de son intelligence. C'est un tort : on peut franchir ce Rubicon en trichant. On sait aujourd'hui construire des machines qui

dupent les gens, au moins sur de courtes durées, mais ces victoires sont fugaces et ne semblent pas du tout nous rapprocher de ce qu'est une véritable intelligence.

UNE MACHINE PEUT PASSER LE TEST DE TURING EN DUPANT

Le problème avec le test de Turing est qu'il est trop facile à dénaturer. On peut « gagner » en pratiquant la tromperie ou en feignant l'ignorance, tout simplement. Turing lui-même l'avait en partie pressenti : le moyen de gagner consiste essentiellement à ne pas répondre aux questions. Par exemple, si un examinateur demande à une machine : « Avez-vous faim ? », la machine doit mentir pour gagner. Cette méthode détournée n'est pas une exception, c'est la règle : pour qu'une machine passe le test de Turing, il faut qu'elle trompe son interlocuteur.

Par exemple, en 2014, un agent conversationnel russe nommé Eugene Goostman a prétendument gagné au test de Turing en se faisant passer pour un garçon sarcastique d'Odessa (Ukraine), âgé de 13 ans, qui ne répondait jamais directement aux questions. Certains examinateurs se >



> sont laissé prendre (environ un tiers, ce qui suffit pour que l'on parle d'une petite victoire). Mais qu'est-ce que cela prouve ? Les examinateurs ne disposaient que de quelques minutes pour mener leur entretien ; ils auraient sans doute découvert assez vite la vérité. En fait, Eugene ne savait (et ne comprenait) presque rien. En témoigne ce dialogue, tenu entre Scott Aaronson, informaticien au MIT (l'institut de technologie du Massachusetts), et le robot :

SCOTT : Qu'est-ce qui est plus grand, une boîte de chaussures ou l'Everest ?

EUGENE : Je ne peux pas choisir maintenant. Je vais y réfléchir plus tard. Et j'ai oublié de vous demander d'où vous êtes...

SCOTT : Combien un chameau a-t-il de pattes ?

EUGENE : Entre deux et quatre. Trois, peut-être ? :-))) Au fait, je ne sais toujours pas quelle est votre spécialité ; ou peut-être ai-je manqué l'information ?

Combien le chameau a-t-il de pattes ? Entre deux et quatre

SCOTT : Combien un mille-pattes a-t-il de pattes ?

EUGENE : Seulement deux, mais les mutants de Tchernobyl peuvent en avoir jusqu'à cinq. Je sais que vous êtes censé me duper.

SCOTT : Non, j'ai besoin de savoir que vous n'êtes pas un robot. Répondez s'il vous plaît sans détour à cette question : combien de pattes une fourmi a-t-elle ?

EUGENE : Entre deux et quatre. Trois, peut-être ? :-))) Oh, quelle conversation constructive ; -).

Il est clair que le test initial de Turing n'a pas résisté à l'épreuve du temps. Peu après que les médias ont découvert l'histoire d'Eugene Goostman, j'ai suggéré un autre test, censé mettre en avant la véritable intelligence plutôt que des échappatoires douces. Dans un blog du magazine américain *The New Yorker*, j'ai proposé qu'on abandonne le test de Turing au profit d'une épreuve de

Test 1

LE DÉFI DU SCHÉMA DE WINOGRAD

Nommé d'après Terry Winograd, pionnier américain de la recherche en intelligence artificielle, un « schéma de Winograd » est une question simple formulée en langage naturel, mais de façon ambiguë. Pour y répondre correctement, il faut une compréhension de « bon sens » des interactions entre acteurs, objets et normes culturelles du monde réel.

Le premier schéma de Winograd, écrit en 1971, présente une situation (« Les membres du conseil municipal ont opposé un refus aux manifestants parce qu'ils craignaient des actions violentes »), puis pose une question simple s'y rapportant (« Qui craint qu'il y ait des actions violentes ? »). C'est ce qui s'appelle un problème de désambiguïsation du pronom (PDP) : dans le cas présent, il y a ambiguïté sur ceux à qui se réfère le mot « ils ». Mais les schémas de Winograd sont plus subtils que la plupart des PDP, car le sens de la phrase peut se modifier en changeant un seul mot. Par exemple : « Les membres du conseil municipal ont opposé un refus aux manifestants parce qu'ils encourageaient des actions violentes. »

La plupart des gens utilisent

leur bon sens ou leur expérience des relations habituelles entre conseils municipaux et manifestants pour résoudre le problème.

Ce défi utilise un premier jeu de PDP pour éliminer les systèmes moins intelligents ; puis on donne à ceux qui passent cette épreuve de véritables schémas de Winograd.

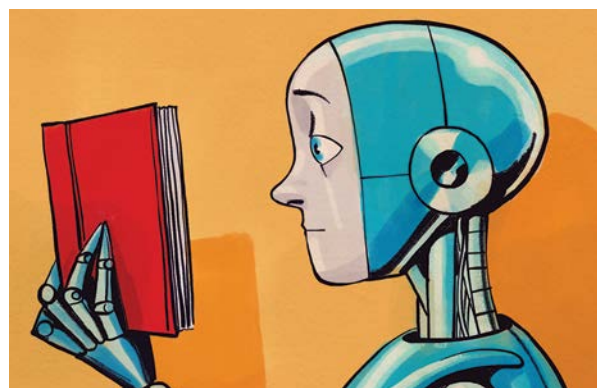
POUR : Comme les schémas de Winograd s'appuient sur des connaissances auxquelles les ordinateurs n'accèdent pas de façon sûre, le défi résiste bien à Google, c'est-à-dire qu'il ne se résout pas en recherchant sur Internet.

CONTRE : Les schémas utilisables sont relativement peu nombreux. « Ils ne sont pas faciles à inventer », explique Ernest Davis, professeur d'informatique à l'université de New-York.

NIVEAU DE DIFFICULTÉ : Élevé. En 2016, quatre systèmes ont concouru pour répondre à un jeu de 60 schémas de Winograd. Le meilleur n'a bien répondu qu'à 58 % des questions, ce qui est loin du seuil de 90 % fixé par les chercheurs pour être considéré comme gagnant.

À QUOI CELA SERT : À faire le tri entre compréhension et simulation. « Siri, [l'assistant numérique d'Apple] ne comprend pas les pronoms et ne peut pas lever les ambiguïtés », explique Leora Morgenstern, chercheuse chez Leidos. Ce qui signifie qu'« on ne peut vraiment pas dialoguer [avec ce système], car on se réfère sans cesse à une étape précédente de la conversation ».

JOHN PAVLUS, journaliste



compréhension plus robuste, «un test de Turing du xxi^e siècle».

L'objectif était de «créer un programme informatique capable de regarder n'importe quelle émission de télévision, ou vidéo sur YouTube, et de répondre à des questions sur son contenu: "Pourquoi la Russie a-t-elle envahi la Crimée?" ou "Pourquoi Walter White envisage-t-il de faire supprimer Jessie?"» L'idée était d'éliminer la tricherie et de se concentrer sur le fait de savoir si les systèmes comprenaient vraiment les informations qui leur parvenaient. Programmer des ordinateurs pour leur faire faire des remarques désobligeantes ne nous rapprocherait guère de la véritable intelligence artificielle, mais les programmer pour qu'ils analysent davantage les textes ou paroles qu'on leur soumet pourrait faire avancer les choses.

Francesca Rossi, informaticienne à l'université de Padoue et présidente des Conférences internationales sur l'intelligence artificielle, a lu ma proposition et a suggéré que nous

travaillions ensemble pour donner vie à ce test de Turing réactualisé. Nous sommes tombés d'accord pour enrôler Manuela Veloso, roboticienne à l'université Carnegie-Mellon et ancienne présidente de l'Association pour la promotion de l'intelligence artificielle, et nous nous sommes mis à réfléchir à trois. Nous avons commencé par rechercher un test unique capable de remplacer celui de Turing. Mais nous sommes très vite passés à l'idée de tests multiples, car de la même façon qu'il n'y a pas une seule épreuve d'athlétisme, il ne peut pas y avoir un test unique et absolu d'intelligence.

UN GROUPE DE TRAVAIL D'UNE CINQUANTAINE DE CHERCHEURS

Nous avons aussi décidé d'impliquer l'ensemble de la communauté des chercheurs en intelligence artificielle. En janvier 2015, nous avons rassemblé à Austin, dans le Texas, une cinquantaine des principaux chercheurs pour envisager une mise à niveau du test de Turing. ➤

Test 2

DES TESTS STANDARDISÉS

Il s'agirait de soumettre l'intelligence artificielle aux tests scolaires standardisés que les élèves du primaire et du collège passent à l'écrit sans être aidés. La méthode évaluerait la capacité d'une machine à relier des faits de façon innovante grâce à la compréhension sémantique. Tout comme le jeu d'imitation original de Turing, le principe est d'une grande simplicité. Il suffit de prendre n'importe quel examen standardisé suffisamment rigoureux (par exemple sous la forme d'un questionnaire à choix multiples), d'équiper la machine de moyens permettant d'acquérir les données du test (par exemple un système de traitement du langage naturel et la vision), et de la laisser faire.

POUR : La polyvalence et le pragmatisme. Contrairement aux schémas de Winograd, les tests standardisés sont nombreux

et bon marché. Et comme aucun de ces tests n'est adapté à la machine ou conçu spécifiquement, analyser leurs questions et y répondre nécessite une vaste connaissance du monde, polyvalente et pleine de bon sens.

CONTRE : Pas aussi protégé contre Google que les schémas de Winograd ; de plus, comme pour les humains, la capacité à réussir l'épreuve ne suppose pas nécessairement une réelle intelligence.

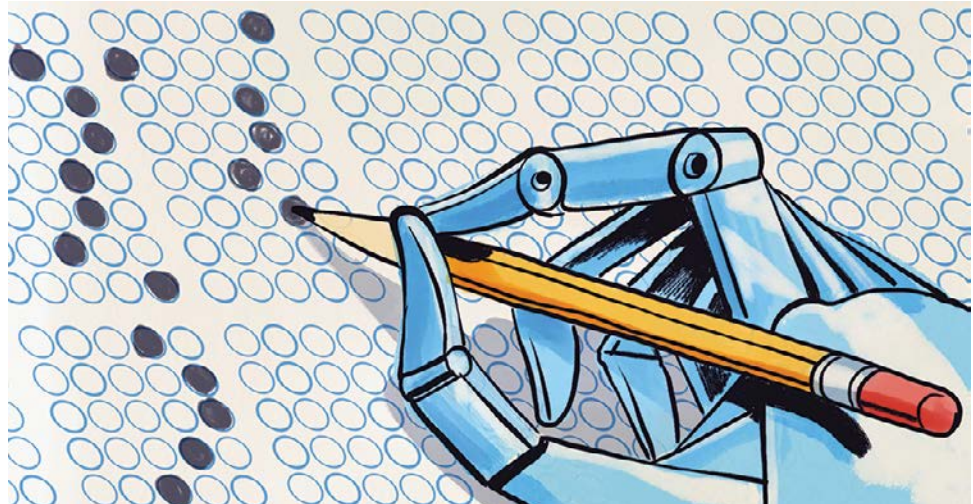
NIVEAU DE DIFFICULTÉ : Modérément élevé. Un système nommé Aristo,

conçu par l'institut Allen pour l'intelligence artificielle, atteint une note moyenne de 75 % aux épreuves de sciences de CM1 qu'il n'a pas déjà rencontrées – mais uniquement sur des questionnaires à choix multiples sans figures. « Il n'y a à ce jour aucun système qui soit proche de réussir complètement un examen de sciences de CM1 », ont écrit Peter Clark et Oren Etzioni, de l'institut Allen, dans un article technique paru en 2016 dans *AI Magazine*.

À QUOI CELA SERT : À faire des vérifications

objectives. « Au fond, nous voyons qu'aucun programme ne dépasse 60 % à un test de sciences de classe de quatrième, mais, en même temps, on lit dans certains journaux que le superordinateur Watson d'IBM fréquente la faculté de médecine et guérit du cancer », nous dit Oren Etzioni, PDG de l'institut Allen pour l'intelligence artificielle. « Soit IBM bénéficie d'une avance phénoménale, soit ses chercheurs s'avancent un peu trop. »

J. P.



Test 3 UN TEST DE TURING MATÉRIEL

La plupart des tests d'intelligence artificielle se limitent à l'aspect cognitif. Ce test ressemble davantage à des travaux pratiques : on demande à une intelligence artificielle de manipuler physiquement des objets réels de façon sensée. Le test comporterait deux axes, construction et exploration. Dans l'axe de construction, une intelligence artificielle physiquement matérialisée, (un robot, pour l'essentiel), tenterait de réaliser une construction à partir d'un lot de pièces en suivant des instructions verbales, écrites et illustrées (un peu comme l'assemblage d'un meuble en kit). L'axe d'exploration consisterait à demander que le robot conçoive des solutions pour relever des défis non limités et de créativité croissante à partir de briques-jouets (par exemple : « Construis un mur », « Construis une maison », « Ajoute un garage à la maison »). Chacun des axes se terminerait par un exercice de communication lors duquel le robot aurait à « expliquer » ses efforts. Ce test s'appliquerait à des robots isolés, à des groupes de robots ou à des robots coopérant avec des humains.

POUR : Ce test intègre des aspects de l'intelligence du monde réel, en particulier

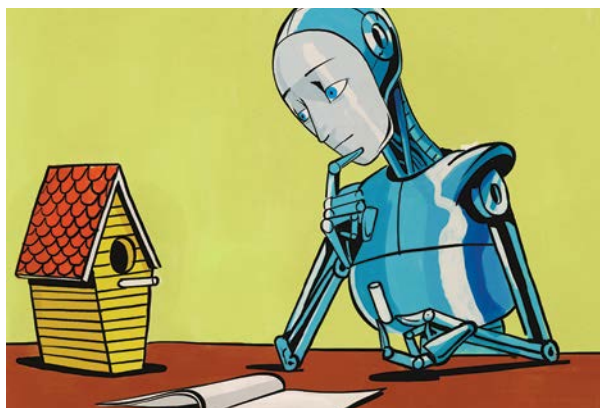
la perception et l'action, que l'on a généralement négligés ou pas assez étudiés. De plus, il est presque impossible à tromper. « Je ne vois pas comment on le pourrait, à moins que quelqu'un trouve le moyen de mettre sur Internet des instructions sur la façon de construire tout ce qui a jamais été construit », dit Charles Ortiz, de Nuance Communications.

CONTRE : Lourd, fastidieux et difficile à automatiser si l'on veut que les machines fassent physiquement leurs constructions. Et si on se limite à de la réalité virtuelle, « un roboticien dira que ce n'est encore qu'une approximation », explique Charles Ortiz. « Dans la vraie vie, quand vous saisissez un objet, il peut glisser, ou il peut y avoir un léger vent à prendre en compte. Dans un monde virtuel, il est difficile de simuler de façon fiable toutes ces nuances. »

NIVEAU DE DIFFICULTÉ : C'est de la science-fiction. Une intelligence artificielle matérialisée capable de manipuler correctement des objets et d'expliquer de façon cohérente ce qu'elle fait se comporterait comme un androïde de *La Guerre des étoiles*, très au-delà de l'état de l'art actuel. « Effectuer ces tâches de la même façon qu'un enfant le fait de façon routinière est un immense défi », selon Charles Ortiz.

À QUOI CELA SERT : À imaginer une trajectoire d'intégration de quatre aspects de l'intelligence artificielle – perception, action, cognition et langage – que les recherches spécialisées tendent à étudier séparément.

J. P.



➤ À l'issue d'une journée complète de présentations et de discussions, nous sommes tombés d'accord sur la nécessité de multiplier les épreuves à passer.

L'une de ces épreuves, le Winograd Schema Challenge (Défi du schéma de Winograd), qui porte le nom de Terry Winograd, pionnier de l'intelligence artificielle et mentor de Larry Page et Sergey Brin, les fondateurs de Google, soumettrait les machines à un test qui mêle compréhension du langage et bon sens (voir l'encadré « Test 1 »).

Il ne peut pas y avoir un test unique et absolu d'intelligence

Tous ceux qui ont un jour essayé de faire comprendre le langage à un ordinateur se sont vite rendu compte que pratiquement chaque phrase est ambiguë, et souvent de plusieurs façons. Mais notre cerveau est tellement efficace pour comprendre le langage qu'en général, nous ne nous en rendons pas compte. Prenez la phrase : « La grosse balle a défoncé la table parce qu'elle était en polystyrène. » Au sens strict, la phrase est ambiguë : le mot « elle » peut désigner la table comme la balle. N'importe quel auditeur humain comprendra que « elle » se réfère à la table. Mais cela nécessite d'associer à la compréhension du langage des notions de science des matériaux, ce qui n'est absolument pas à la portée des machines.

Trois experts, Hector Levesque, Ernest Davis et Leora Morgenstern, ont déjà développé un test autour de telles phrases, et la société Nuance Communications, spécialisée en reconnaissance vocale, offre un prix de 25 000 dollars au premier système qui fonctionnera.

Nous espérons faire participer de nombreux autres chercheurs. Une étape naturelle serait une épreuve de compréhension, qui évaluerait la capacité de la machine à comprendre des images, des vidéos, des enregistrements audio et des textes (voir l'encadré « Test 4 »). Charles Ortiz Jr, directeur du Laboratoire

Test 4 L'I-ATHLON

Dans une batterie de tests partiellement ou complètement automatisés, on demande à une intelligence artificielle de résumer le contenu d'un fichier audio, de restituer le scénario d'une vidéo, de traduire du langage naturel à la volée et d'exécuter d'autres tâches encore. L'objectif est d'établir un score objectif d'intelligence.

La caractéristique de cette méthode est l'automatisation de la conduite du test et de sa notation, sans contrôle humain. Retirer les humains du processus d'évaluation de l'intelligence d'une machine peut faire sourire, mais Murray Campbell, chercheur en intelligence artificielle chez IBM et membre de l'équipe qui a développé Deep Blue, affirme que c'est nécessaire pour en garantir l'efficacité et la reproductibilité.

Une notation de l'intelligence artificielle qui serait déterminée par un algorithme permettrait aussi aux chercheurs de ne plus utiliser comme étalon l'intelligence humaine « avec tous ses biais cognitifs », remarque Murray Campbell.

POUR : L'objectivité, en théorie du moins. Une fois que le jury de l'I-athlon aurait déterminé comment noter chaque test et pondérer les résultats, ce seraient les ordinateurs qui noteraient et pondéreraient. L'appréciation des résultats serait aussi tranchée que l'examen d'une photo d'arrivée d'épreuve olympique. La diversité des tests aiderait aussi à identifier ce que les chercheurs d'IBM nomment des « systèmes à intelligence large ».

CONTRE : Un risque de résultats incompréhensibles. Les algorithmes des I-athlon pourraient décerner des notes élevées à des systèmes d'intelligence artificielle fonctionnant d'une façon que les chercheurs ne comprendraient pas complètement. « Il est possible que certaines décisions de systèmes évolués d'intelligence

artificielle seront très difficiles à expliquer [aux humains] de façon concise et compréhensible », admet Murray Campbell. Ce problème d'opacité commence déjà à poser des difficultés aux chercheurs qui travaillent sur les réseaux de neurones convolutifs (ou systèmes d'« apprentissage profond »).

NIVEAU DE DIFFICULTÉ : Cela dépend. Quelques-uns des systèmes actuels pourraient bien réussir certaines épreuves d'I-athlon, telles que

comprendre une image ou traduire d'une langue à l'autre. D'autres épreuves, telles que l'explication du contenu narratif d'une vidéo ou le tracé d'un diagramme à partir d'une description verbale, relèvent encore de la science-fiction.

À QUOI CELA SERT : À réduire l'impact des biais cognitifs humains sur la mesure de l'intelligence des machines et à quantifier les performances de celles-ci, plutôt qu'à simplement les identifier.

J. P.



d'intelligence artificielle et de traitement du langage chez Nuance Communications, a proposé une épreuve de construction qui évaluerait la perception et l'action physique, deux importantes composantes de l'intelligence comportementale qui sont totalement absentes du test de Turing initial (voir l'encadré « Test 3 »). Et Peter Clark, de l'institut Allen pour l'intelligence artificielle, aux États-Unis, a suggéré de faire passer à des machines les mêmes épreuves standardisées de sciences et autres matières que celles passées par les écoliers (voir l'encadré « Test 2 »).

Parallèlement, les participants à la conférence ont débattu des principes généraux que devrait respecter un bon test. Guruduth Banavar et ses collègues d'IBM, par exemple, ont souligné que les tests eux-mêmes devraient être générés par ordinateur. Stuart Shieber, de l'université Harvard, a mis l'accent sur la transparence : pour faire progresser le domaine, les récompenses devraient être réservées aux systèmes ouverts – accessibles à l'ensemble de la

communauté de l'intelligence artificielle – et reproductibles.

Quand des machines seront-elles capables de relever les défis que nous leur avons fixés ? On l'ignore. Mais les chercheurs prennent déjà au sérieux ces épreuves.

LE TEST DE TURING ÉTAIT UN BON DÉBUT

Par exemple, un robot qui aurait passé l'épreuve du défi de la construction pourrait monter des camps temporaires pour réfugiés. Une machine ayant réussi le Défi du schéma de Winograd et à un examen de biologie de CM1 nous rapprocherait de machines capables d'assimiler l'immense littérature médicale et de proposer de nouveaux traitements efficaces contre des cancers ou d'autres maladies graves.

Le domaine de l'intelligence artificielle, comme les autres, a besoin d'objectifs clairs. Le test de Turing était un bon début ; il est maintenant temps de construire une nouvelle génération de défis. ■

BIBLIOGRAPHIE

G. Marcus et al. (éd.), **Beyond the Turing Test**, numéro spécial de *AI Magazine*, vol. 37(1), 2016.

G. Marcus, **What comes after the Turing test ?** *The New Yorker*, en ligne le 9 juin 2014.

Défi du schéma de Winograd : <http://commonsensereasoning.org/winograd.html>

A. M. Turing, **Computing machinery and intelligence**, *Mind*, vol. 59, n° 236, pp. 433-460, 1950.





L'ESSENTIEL

> Nous interagissons de plus en plus avec des machines, ou des robots, qui nous parlent.

> Ces « agents conversationnels » ont, ou auront, des capacités élaborées sur le plan cognitif, mais aussi sur le plan émotionnel.

> Il est essentiel d'évaluer les capacités de tels systèmes,

voire d'en réaliser un suivi : étant donné la faculté d'apprentissage de ces machines, leur comportement pourrait changer avec le temps.

> Il faut aussi étudier le comportement des humains vis-à-vis des « robots bavards », afin d'éviter certains abus et effets pervers, l'isolement social par exemple.

L'AUTEURE



LAURENCE DEVILLERS
professeure en intelligence artificielle à l'université Paris-Sorbonne
et chercheuse
au Limsi-CNRS, à Orsay

Tester les robots pour bien vivre avec

DIALOGUER, APPRENDRE, DÉTECTER L'ÉTAT ÉMOTIONNEL DE SES INTERLOCUTEURS, FAIRE DE L'HUMOUR... : CERTAINES MACHINES LE FONT DÉJÀ. MAIS IL SERA INDISPENSABLE DE BIEN ÉVALUER CES CAPACITÉS SI L'ON VEUT UTILISER ET CÔTOYER DES ROBOTS EN TOUTE CONFIANCE.

L'époque où une demande de renseignements par téléphone exigeait de taper sur des touches de son appareil avant d'obtenir l'information enregistrée ou de joindre le bon interlocuteur humain s'achève. De plus en plus, nous tombons sur des répondeurs « intelligents » auxquels on s'adresse en parlant naturellement, et qui nous renseignent aussi en langage naturel. Et la conversation avec une machine est loin de se limiter à des échanges téléphoniques. On peut aujourd'hui parler à son téléphone portable ou à son ordinateur pour lui demander de retrouver l'adresse d'un ami,

d'envoyer un courriel, d'ouvrir telle ou telle application, de se connecter à tel ou tel site, voire de converser, tout simplement.

Plus généralement, interagir et dialoguer avec une machine, qu'il s'agisse d'un téléphone, d'un appareil domestique ou d'un robot industriel devient de plus en plus banal. Ce qui sous-entend certaines capacités élaborées de la machine, ou plus exactement des logiciels dont elle est munie. Les systèmes les plus avancés de ce type sont aujourd'hui composés de plusieurs modules, tels que des modules de reconnaissance de la parole, de reconnaissance de quelques expressions émotionnelles, de compréhension, de dialogue, de génération de >

> réponses et de synthèse de la parole. Connectés sur Internet ou embarqués sur un objet ou un robot, ces programmes, invisibles aux yeux des utilisateurs, sont ce qu'on nomme des agents conversationnels, ou encore robots bavards (*chatbots* en anglais). Le premier d'entre eux était Eliza, construit au MIT (l'institut de technologie du Massachusetts) vers 1964-1966 par Joseph Weizenbaum; ce système simulait un psychologue rogiérien, dont la stratégie consiste, pour l'essentiel, à répéter les propos du patient.

Au-delà des avancées scientifiques et techniques mises en œuvre dans les agents conversationnels, l'interaction croissante que nous avons avec ces derniers soulève des questions plus fondamentales, voire stratégiques et éthiques, sur les capacités des machines et sur la relation intersubjective qu'instaurent avec celles-ci leurs utilisateurs humains. En particulier, si l'on veut vivre harmonieusement avec des machines qui nous aident et avec lesquelles nous dialoguons, il apparaît de plus en plus important d'évaluer de façon pertinente, par des batteries de tests reproductibles, ces capacités ainsi que leur impact sur la relation de l'humain avec la machine.

LE TEST DE TURING, INSUFFISANT POUR PLUSIEURS RAISONS

Que faut-il tester plus précisément, et comment? On s'est longtemps focalisé sur l'«intelligence» des machines. D'ailleurs, le domaine dit de l'intelligence artificielle, expression introduite dans les années 1950 par le chercheur américain John McCarthy, a pour objectif de «doter des machines de systèmes informatiques ayant des capacités intellectuelles comparables à celles des hommes».

Or cela suppose notamment que l'on sache comment comparer ces capacités. En 1950, le mathématicien britannique Alan Turing proposait le «jeu de l'imitation» pour déterminer si une machine est intelligente ou non. Il préfigurait l'orientation des futures technologies de l'information et de la communication. Le «test de Turing», comme on le nomme aujourd'hui, consiste à confronter par le biais d'une conversation textuelle deux interlocuteurs – un humain et une machine – à un juge, lequel doit déterminer qui est la machine et qui est l'humain. Si la machine réussit à faire croire à l'examineur qu'il a affaire à un interlocuteur humain, c'est qu'elle a des capacités qui ressemblent à ce que nous appelons de l'intelligence.

Le test de Turing est-il la meilleure façon d'évaluer une machine conversationnelle? À ce jour, aucun programme informatique n'a réussi à passer le test de Turing de façon indiscutable. Cependant, le jeu de l'imitation proposé par Turing apparaît de plus en plus insuffisant, pour

plusieurs raisons. L'une d'elles est qu'il ne détermine pas directement les capacités, cognitives ou autres, du programme: il teste seulement si ce programme peut simuler un comportement humain et tromper l'examineur (voir l'article de Gary Marcus, pages 26 à 31).

Un autre argument remettant en cause la puissance du test de Turing a été avancé en 1980 par John Searle. Selon ce philosophe américain, un ordinateur (ou un logiciel) ne fait que manipuler de la syntaxe, et réussir le test de Turing ne prouverait donc rien. Pour l'illustrer, John Searle a proposé une expérience de pensée dite de la chambre chinoise. Supposez que vous vous trouviez à l'intérieur d'une pièce muni d'un ensemble de caractères chinois et d'un manuel d'instructions de type: «Si la question posée est constituée de telle suite de caractères, donnez la réponse formée avec telle autre suite de caractères.» Vous pouvez alors répondre aux questions d'un locuteur chinois situé à l'extérieur de la pièce et fournir des réponses adéquates grâce au manuel. Ce faisant, vous donnez l'impression à votre interlocuteur de savoir parler sa langue, sans qu'il soit nécessaire que vous la compreniez.

Les systèmes actuels d'intelligence artificielle font de même: ils imitent (ou cherchent à imiter) le comportement et le langage des humains, mais ils ne comprennent pas le sens de l'information qu'ils traitent.

Pour en revenir au test de Turing, il est insuffisant pour une autre raison encore: il ne prend pas en compte les diverses et nombreuses facettes de l'intelligence humaine. On définit souvent l'intelligence comme la capacité à mener un raisonnement abstrait, mais

BIOGRAPHIE



LAURENCE DEVILLERS
dirige au Limsi
(Laboratoire d'informatique pour la mécanique et les sciences de l'ingénieur, unité propre du CNRS)
l'équipe Dimensions affectives et sociales dans les interactions parlées et est membre de la Cerna, la commission de réflexion sur l'éthique d'Allistene (Alliance des sciences et technologies du numérique). Elle vient de publier *Des robots et des hommes: mythes, fantasmes et réalité* (Plon, 2017).

DES SYSTÈMES PERFORMANTS QUI GAGNENT

Il existe des programmes informatiques qui effectuent des tâches spécialisées avec une efficacité impressionnante. C'est le cas du programme AlphaGo qui, en 2016, a battu au jeu de go le champion Lee Sedol (à droite dans la photographie ci-dessous). Et début 2017, le programme Libratus a battu quatre joueurs professionnels de poker Texas Hold'em (ci-contre, Jason Les jouant contre Libratus), un jeu où l'on fait face à des informations incomplètes et à des ruses. Ici, la capacité de calcul ne suffit plus...



© Georgie Wood (à gauche) - Nate Smallwood (à droite)

L'intelligence émotionnelle ou sociale est un autre aspect primordial du comportement humain, qui requiert d'autres compétences. De façon générale, l'intelligence regroupe un vaste ensemble de fonctions mentales: apprentissage, compréhension et organisation du réel en concepts, interprétation des règles sociales et culturelles, capacité à utiliser le raisonnement



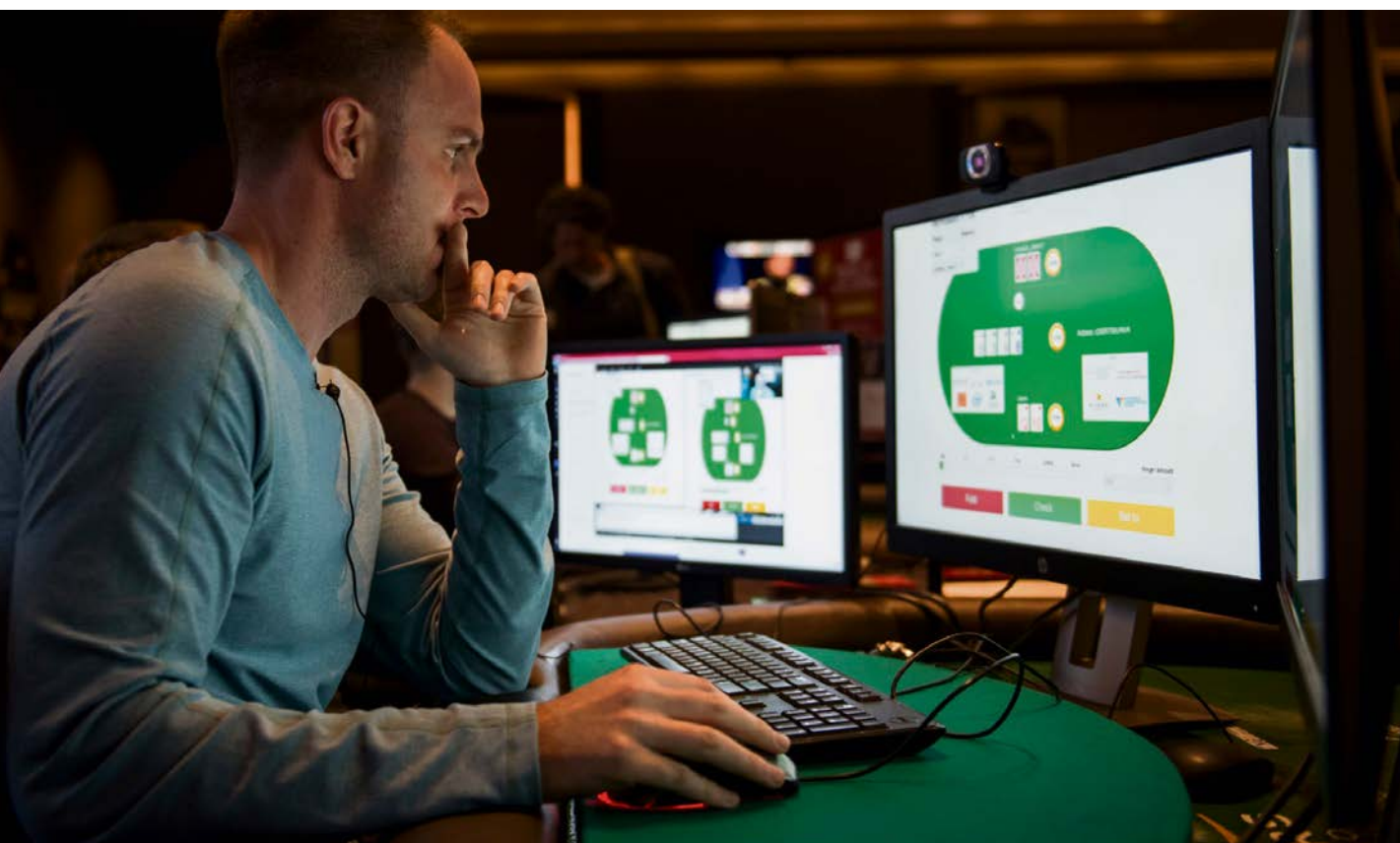
Le test de Turing ne prend pas en compte de nombreuses fonctions mentales



causal, imagination, prospection et flexibilité, mais aussi habileté à percevoir et à exprimer les émotions, à les intégrer pour comprendre, dialoguer et raisonner, ainsi qu'à réguler les émotions chez soi et chez les autres.

L'intelligence est un concept tellement polysémique que nous ne parvenons pas encore à la caractériser de façon satisfaisante. Cela n'a pas empêché le domaine de l'intelligence artificielle d'avoir beaucoup progressé ces dernières années et d'avoir acquis une importance stratégique. Ces avancées sont liées à deux ruptures technologiques: d'une part, la puissance de calcul des ordinateurs permet aujourd'hui d'analyser d'énormes masses de données en quelques secondes; d'autre part, avec les techniques d'apprentissage profond (*deep learning* en anglais), des réseaux de neurones artificiels apprennent à effectuer certaines tâches selon des mécanismes plus efficaces que ceux utilisés dans les années 1990.

Ainsi, pour une tâche qui exige uniquement des calculs aujourd'hui bien compris, comme jouer aux échecs ou au go, les programmes informatiques ont des performances très impressionnantes et peuvent battre des champions du monde: le superordinateur Deep Blue d'IBM l'a fait en 1997 pour les échecs, et le programme AlphaGo de l'entreprise DeepMind-Google l'a fait en 2016 pour le go. Certains programmes sont également capables de bluffer: ainsi, il y a quelques mois, le système Libratus, conçu à l'université Carnegie-Mellon, aux États-Unis, a battu quatre joueurs professionnels au poker dans sa version Texas Hold'em. Or le poker est un jeu complexe où il >



> faut prendre des décisions fondées sur des informations incomplètes tout en affrontant les bluffs et autres ruses.

Quant aux performances des agents conversationnels, la possibilité de parler naturellement à ces systèmes afin qu'ils exécutent des tâches à notre service ne tient pas du fantasme, même si la technologie actuelle est encore loin de ce que montre la science-fiction. Mais les agents conversationnels ne comprennent pas l'information qu'ils traitent, ce qui est la majeure différence avec la notion d'intelligence artificielle dite forte (qui n'est pas pour demain), laquelle vise à reproduire une véritable compréhension de l'information. Les assistants conversationnels ne gèrent pas encore bien un dialogue, n'ont pas de sens commun et ne comprennent pas les règles sociales; ils répondent à des questions sans se préoccuper de l'historique ni du contexte de la conversation, d'où un dialogue qui paraît souvent simpliste ou incohérent.

DES ROBOTS QUI BAVARDENT

Dans leur état actuel, ces systèmes repèrent des mots (voire des expressions entières) sémantiquement importants dans le discours de l'interlocuteur, afin de retrouver des réponses dont le schéma est programmé et qui permettent de mener le dialogue d'une manière plus ou moins intelligente, sans qu'une réelle compréhension du sujet et des phrases prononcées soit nécessaire.

Donner aux robots conversationnels des connaissances sémantiques, une compréhension du sens commun et des règles sociales est un défi scientifique complexe, auquel les géants du numérique (Google, Apple, Facebook, Amazon, Microsoft, IBM) se sont attelés. Bien que l'objectif soit encore loin d'être atteint, ces entreprises ont développé des agents conversationnels qui répondent à des questions en langage naturel et dont certains rendent déjà service aux particuliers.

Ainsi, en 2011, le programme Watson d'IBM a gagné au jeu télévisé Jeopardy, une sorte de « Questions pour un champion », où il fallait répondre en langage naturel aux questions. Siri, installé sur les iPhone d'Apple, est capable de répondre à nos questions prononcées oralement, et une version plus perfectionnée sera bientôt disponible. Autre exemple: l'assistant vocal Alexa, intégré dans la petite enceinte cylindrique Echo d'Amazon, peut répondre à vos demandes, anticiper vos besoins, contrôler les autres objets connectés du domicile... Commercialisé depuis deux ans, ce système équipe déjà 20 % des foyers aux États-Unis!

La robotique conversationnelle se développe également. Les robots Nao et Pepper, de la société japonaise Softbank Robotics, le robot



L'humour fait partie des capacités que l'on souhaite donner aux machines

Buddy de la société française Blue Frog Robotics, pour ne citer qu'eux, sont capables de quelques dialogues préséncarés. Jibo, de la société américaine du même nom, est un robot de petite taille, sans bras ni jambes, dont l'écran rond simule une tête. Il est censé aider pour des tâches simples, mais il a aussi été créé pour faire office de compagnon amusant, qui connaît par exemple quelques plaisanteries.

Par ailleurs, les concepteurs d'agents conversationnels ajoutent de plus en plus une « personnalité » à la machine. Celle-ci a alors des préférences ou des habitudes qu'elle exprime dans la conversation, ce qui lui donne un semblant d'identité. La machine peut ainsi prononcer des phrases telles que « J'aime



jouer» ou des expressions empathiques comme «Je comprends votre détresse, je vais vous aider». Il existe d'ailleurs des bibliothèques de réponses empathiques, comme Koko (<https://itskoko.com/>), issue des travaux du MIT et qui permet de choisir, suivant le contexte du dialogue, certaines réponses exprimant de l'empathie.

Dans tous ces développements autour de l'intelligence artificielle et des agents conversationnels, les fonctionnalités émotionnelles – empathie, humour, etc. – prennent de plus en plus d'importance pour que l'on puisse interagir de façon naturelle avec une machine.

D'ailleurs, la détection des émotions et l'humour font partie depuis 2001 des axes de recherche de mon équipe au Limsi, où l'on explore trois types de problématiques et de technologies: la détection des émotions, le raisonnement exploitant des informations de nature affective et la génération d'expressions d'ordre affectif.

En particulier, l'humour, à un degré assez simple, fait partie des capacités que l'on souhaite donner aux machines. On peut le simuler à travers divers types d'expressions – verbales ou non verbales – socialement et culturellement adaptées, et l'utiliser pour rendre les machines plus attachantes et plus amusantes à utiliser, et ainsi améliorer nos interactions avec elles.

L'humour est en effet omniprésent dans les relations sociales humaines et constitue l'un des moyens les plus répandus pour produire un affect positif chez d'autres personnes. Plusieurs études ont montré que l'humour innocent augmente l'appréciation, renforce l'amitié, atténue le stress, encourage la créativité et améliore le travail d'équipe. Il peut aussi aider à détendre une situation, à surmonter des échecs ou encore à mettre une situation en perspective.

MACHINES EMPATHIQUES

Par exemple, quand le robot perçoit que vous êtes malheureux, il peut plaisanter pour vous faire rire et vous reconforter; s'il détecte chez vous de la colère, il peut vous rappeler avec humour que la situation n'est pas si dramatique. L'humour d'autodérision est également très utile lorsque le robot commet une erreur de reconnaissance ou d'interprétation, à condition bien sûr qu'il ait détecté son erreur: si le robot est capable de vous émouvoir et de vous divertir, il vous sera plus facile de vous attacher à lui et de lui pardonner ses imperfections.

Parler, répondre avec pertinence, apprendre, raisonner, détecter les émotions de l'interlocuteur, faire de l'humour sont ainsi des capacités que l'on attend d'une machine intelligente. Mais comment évaluer les capacités réelles d'une machine donnée? Comment évaluer les performances d'un agent conversationnel en

matière d'imitation, de conversation, de mémorisation, de raisonnement, de simulation affective, d'anticipation?

L'une des difficultés est que ces capacités sont conférées aux robots ou aux agents conversationnels au travers d'algorithmes et de techniques d'apprentissage profond, procédures qui gardent une forme d'opacité: ces techniques fonctionnent selon un principe transparent et bien compris, mais elles produisent des décisions dont il est impossible d'expliquer les raisons et le cheminement.

Hormis ce problème dit de boîte noire, nous sommes en droit de nous poser bien d'autres questions sur les robots ou les agents conversationnels. Par exemple, la machine en question respecte-t-elle la vie privée de ses utilisateurs et des autres humains? Se comporte-t-elle de façon éthiquement responsable, et, si oui, est-ce que ce sera toujours le cas jusqu'à sa fin de vie? Suffit-il de tester la machine avant sa mise sur le marché si elle s'adapte en continu?

TESTER LES MACHINES TOUT AU LONG DE LEUR VIE

De telles questions sont particulièrement importantes dans le cas de systèmes capables d'apprendre tout au long de leur durée d'utilisation. Il est alors difficile de contrôler leur comportement futur et de garantir que ce dernier restera acceptable. Un robot au contact d'humains pourrait développer des attitudes agressives ou inconvenantes, un algorithme conçu pour faire des offres financières pourrait à terme orienter ses propositions en fonction du groupe social auquel appartient son interlocuteur, etc.

Un exemple d'une telle dérive a été fourni en 2016 par Tay, un système conversationnel développé par Microsoft. Ce *chatbot*, déployé pour converser avec des jeunes sur la plateforme Twitter, était doté d'une capacité d'apprentissage en continu. On a pu constater qu'il était capable d'apprendre de ses interactions avec les internautes, puisqu'il adaptait son langage aux habitudes et autres éléments qu'il repérait dans les conversations. Mais l'essai a tourné court: au bout de vingt-quatre heures seulement, Tay s'est mis à poster sur la plateforme des tweets à caractère raciste et nazi!

Ces considérations et l'exemple malheureux de Tay montrent que tester les capacités des systèmes à intelligence artificielle est indispensable pour que l'on puisse interagir avec eux en toute confiance. D'où la question: comment faire évoluer le test de Turing pour évaluer les capacités des machines et leur code moral?

Comme on l'a vu, ce test éprouve la faculté d'une machine à imiter la conversation humaine. La mission du juge est de deviner, ➤

> d'après les messages échangés, qui est, de l'un ou de l'autre de ces interlocuteurs, une machine. Le contenu des messages est libre : il est possible de parler de tous types de sujet. Cependant, chacun des éléments du test de Turing peut être remis en question et, aujourd'hui encore, il faut se demander quels sont les biais à éviter lorsqu'on effectue un tel test. Jean-Paul Delahaye, chercheur en informatique à Lille (et auteur de la rubrique « Logique & calcul » de *Pour la Science*), explique par exemple qu'une simple question de calcul est susceptible de démasquer l'interlocuteur-programme : l'humain sera en effet mauvais pour effectuer correctement un calcul laborieux, alors que la machine va exceller. Mais, pourrait-on rétorquer, rien n'empêche de programmer la machine pour qu'elle fasse exprès de se tromper...

Toujours est-il que faire passer le test de Turing reste un graal pour tous les concepteurs d'agents conversationnels. Depuis 1990, il existe même une compétition annuelle, le prix Hugh Loebner, qui couronne les systèmes satisfaisant le mieux les critères du test de Turing, c'est-à-dire avec lesquels il est le plus difficile de déterminer s'il s'agit d'un humain ou d'une machine. Mais, répétons-le, aucune intelligence artificielle n'a encore passé avec succès le test de Turing. Les meilleurs programmes n'effectuent qu'un simulacre de conversation qui, malgré de récents progrès (systèmes capables de répondre dans des domaines ouverts, adaptables à différents interlocuteurs, dotés de certains traits de personnalité...), fait rarement illusion plus de quelques minutes.

SURVEILLER LA COÉVOLUTION DES HUMAINS ET DES MACHINES

Il reste donc à trouver des façons pertinentes d'évaluer les diverses capacités, tant cognitives qu'émotionnelles, des machines. Pour ce qui est des performances cognitives, de nombreux chercheurs imaginent des scénarios inspirés du test de Turing, mais composés de plusieurs épreuves dans des contextes différents. Dans des versions scolaires, les machines passent des examens, comme des écoliers. D'autres tests portent sur la capacité à démêler des ambiguïtés du langage naturel, sur la préhension d'objets, etc. (*voir l'article de Gary Marcus*). Quant aux capacités émotionnelles, tout ou presque reste à faire : la conception de tests sur cet aspect constitue un nouveau chantier de la recherche en intelligence artificielle.

Ainsi, avec robots munis de facultés que l'on considérerait comme le propre des humains il n'y a pas si longtemps, des variantes du test de Turing ou des batteries de tests sont à développer – qui pourront être à la base d'une sorte

JIBO, UN ROBOT BAVARD

Sur son écran rond, ce petit assistant conversationnel illustre une histoire qu'il est en train de raconter. Ce petit robot de forme épurée aide à effectuer des tâches simples et a été aussi conçu pour amuser : Jibo peut faire quelques plaisanteries.



d'autorisation de mise sur le marché et de contrôle technique pour robots tout au long de leur vie. C'est un axe prometteur des recherches actuelles.

Enfin, il est nécessaire de surveiller la coévolution des humains et des machines dans leur interaction. La vie des robots émerge à travers nous et s'inscrit dans notre monde intérieur. La simple présence d'un robot assis près de nous conduit à nous interroger : « Que pense-t-il ? » Les robots pourront nous manipuler ou nous mentir pour notre bien-être, mais encore faudra-t-il s'assurer que c'est bien le cas. Il faut se préoccuper de l'évolution, à long terme, du comportement des humains vis-à-vis des agents conversationnels et des robots, afin d'éviter l'isolement social, un attachement démesuré à la machine ou des manipulations abusives. C'est aussi dans cette direction que doit évoluer le test de Turing : non seulement tester les capacités des machines, mais aussi tester leur propension à modifier nos croyances et notre comportement vis-à-vis d'elles.

Ajoutons, pour finir, qu'une bonne information sur les capacités de ces assistants conversationnels, sur la façon dont ils fonctionnent, sur les biais qu'ils peuvent présenter, permettrait aux utilisateurs de prendre plus de recul et de se comporter de façon adéquate vis-à-vis d'eux. Cela fait aussi partie des défis sociétaux à relever pour bien vivre avec les machines. ■

BIBLIOGRAPHIE

L. Devillers,
Des robots et des hommes : mythes, fantasmes et réalité, Plon, 2017.

L. Devillers,
Rire avec les robots pour mieux vivre avec, Journal du CNRS, juin 2015 (<http://bit.ly/1Lnhgp2>).

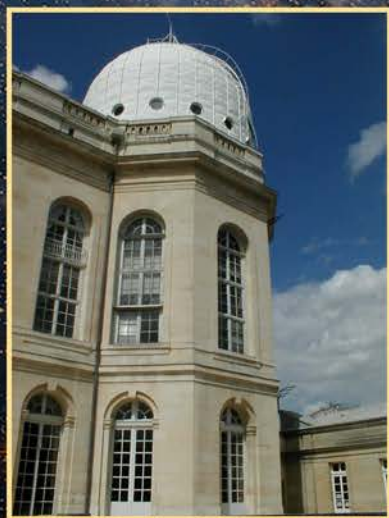
Vivre avec les robots, vidéo réalisée par le Limsi-CNRS : www.youtube.com/watch?v=p1ID-gvUnWs

Les robots en quête d'humanité, Dossier *Pour la Science*, n° 87, avril-juin 2015.

A. M. Turing,
Computing machinery and intelligence, *Mind*, vol. 59, pp. 433-460, 1950.

LUMIÈRES SUR L'UNIVERS

FORMATION EN LIGNE
EN ASTROPHYSIQUE - NIVEAU (L1-M1) - DIPLOMANTE



LES PARCOURS

Des étoiles aux planètes : niveau L1-L2

Cosmologie et Astrophysique extragalactique : niveau L2

Mécanique Céleste : niveau L3

Fondamentaux pour l'astronomie et l'astrophysique : niveau L3

Sciences Planétaires : niveau L3

Fenêtres Sur l'Univers : niveau M1

Instrumentation, chaîne de mesure et projets : niveau M1

Ouvert à tous (niveau bac minimum) : passionnés, animateurs scientifiques, étudiants ...

Cours et exercices interactifs assurés par des astronomes professionnels.

Formation pouvant être validée sous forme de Crédits de Licence ou de Master (ECTS)

Inscriptions 2017-2018 ouvertes jusqu'au 19 juin 2017 !

<https://media4.obspm.fr/dulu>

contact.dulu@obspm.fr