

La Conférence de Dartmouth, naissance de l'Intelligence Artificielle

(©AAAI, AI magazine, vol. 27, 4, 2006)

Le projet de recherches d'été de Dartmouth sur l'intelligence artificielle est un colloque qui est largement considéré comme le moment fondateur de l'intelligence artificielle en tant que champ de recherche. Ce colloque s'est tenu pendant huit semaines à Hanovre, dans le New Hampshire, en 1956, et la conférence a réuni 20 des plus brillants esprits en informatique et sciences cognitives pour un colloque dédié à la conjecture suivante : "tout aspect de l'apprentissage et de n'importe quelle caractéristique de l'intelligence peut être si précisément décrit qu'en principe, une machine devrait pouvoir être fabriquée pour simuler l'intelligence". Des tentatives seront menées pour trouver comment fabriquer des machines étant capables d'utiliser le langage naturel, de formuler des abstractions et des concepts, de résoudre des sortes de problèmes habituellement réservés aux humains, et de s'améliorer elles-mêmes. Nous pensons qu'une avance significative de l'un ou plus de ces problèmes peut avoir lieu si un groupe sélectionné de scientifiques travaillent ensemble sur ce sujet pendant la durée d'un été.

Une proposition pour le projet d'été de recherches sur l'intelligence artificielle (McCarthy et al, 1955)

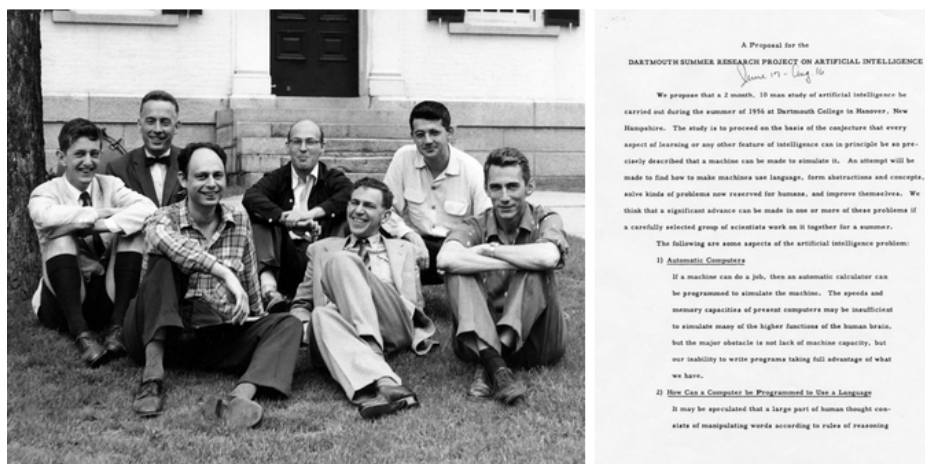


FIGURE 1 — Minsky au centre, Shannon à droite, et d'autres participants à la conférence

Cet article est la traduction d'un article de Jørgen Veisdal, qui a écrit quelques articles en lien avec les mathématiques, consultable dans le paradis de Cantor ici <https://medium.com/cantors-paradise/the-birthplace-of-ai-9ab7d4e5fb00>

Motivation

Avant la conférence, le professeur assistant de mathématiques à Dartmouth John McCarthy et Claude Shannon du MIT avaient co-édité le Volume 34 du journal des Annales d'études mathématiques qui devait paraître dans un avenir proche, sur des études des automates (Shannon & McCarthy, 1956). Les automates sont des machines autonomes destinées à suivre des séquences prédéterminées d'opérations ou à répondre à des instructions prédéterminées. Comme mécanismes d'ingénierie, les automates apparaissent dans une large variété d'applications tels que les horloges mécaniques où un marteau tape sur une cloche toutes les heures et un coq sort pour chanter.

Selon James Moor (2006), McCarthy en particulier avait été déçu par les soumissions d'articles à paraître et leur incapacité à se focaliser sur les possibilités des ordinateurs possédant une intelligence au-delà du plutôt trivial et simple comportement des automates :

A ce moment-là, je pensais : “si nous pouvions obtenir que toute personne intéressée par le sujet lui consacre du temps en évitant la distraction, nous pourrions faire de réelles avancées sur le sujet”. (John McCarthy)

Le groupe initial que McCarthy avait en tête incluait Marvin Minsky qu'il avait connu lorsqu'ils étaient tout deux étudiants ensemble à Fine Hall au début des années 1950. Ces deux-là avaient parlé d'intelligence artificielle et la thèse de Minsky en mathématiques avait comme sujet les réseaux de neurones (Moor, 2006) et la structure du cerveau humain (Nasar, 1998). Ils avaient tous les deux été recrutés aux côtés de Claude Shannon en 1952. Peu de temps avant, McCarthy avait appris que Shannon, qui était bien plus vieux qu'eux, était intéressé par le domaine, également. Finalement, McCarthy s'était rué vers Nathaniel Rochester au MIT quand IBM leur avait offert un ordinateur. Il montra lui aussi de l'intérêt pour l'intelligence artificielle (McCorduck, 1979), et ainsi, les quatre se mirent d'accord pour soumettre à la fondation Rockefeller une proposition pour une conférence / colloque.

La proposition

Nous proposons qu'une étude 2 mois-10 hommes de l'intelligence artificielle soit menée pendant l'été 1956 au Collège de Dartmouth, à Hanovre dans le New Hampshire.

La proposition avait été envoyée par McCarthy et Minsky à la fondation Rockefeller à Dartmouth, puis à Harvard. Ils amenèrent leur proposition aux doyens de la faculté,

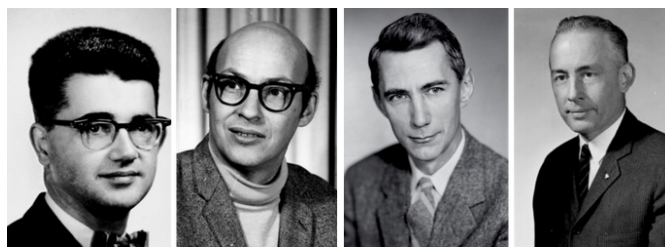


FIGURE 2 — Les initiateurs de la conférence : John McCarthy, Marvin Minsky, Claude Shannon, Nathaniel Rochester)

à Claude Shannon aux Bell Labs et à Nathaniel Rochester à IBM et ils obtinrent leur soutien (Crevier, 1993). La proposition suivit, et fut reproduite dans le magazine AI (McCarthy et al, 1955, Magazine AI Volume 27, Numéro 4 p. 12-14) :

Proposition pour le projet de recherches d'été de Dartmouth sur l'intelligence artificielle (21 août 1955)

1. Machines automatiques :

Si une machine peut réaliser une tâche, alors un calculateur automatique peut être programmé pour simuler la machine. La vitesse et la capacité mémoire des ordinateurs actuels peuvent être insuffisants pour simuler un grand nombre des hautes fonctions intellectuelles du cerveau humain, mais l'obstacle principal n'est pas le manque de capacités de la machine, mais notre incapacité à écrire des programmes qui tirent totalement parti des machines dont nous disposons.

2. Comment un ordinateur peut-il être programmé pour utiliser un certain langage ? On peut supposer qu'une large part de la pensée humaine consiste à manipuler des mots selon des règles de déduction et des règles d'induction. De ce point de vue, former une généralisation consiste à admettre un nouveau mot et quelques règles selon lesquelles des phrases contenant ce mot impliquent et sont impliquées par d'autres. Cette idée n'a jamais été très précisément formulée et n'a pas été testée sur des exemples.

3. Réseaux de neurones

Comment un ensemble de neurones (hypothétiques) peuvent-ils être arrangés de manière à former des concepts. Un travail théorique et expérimental considérable a été réalisé sur ce problème par Uttley, Rashevsky et son groupe, Farley et Clark, Pitts et McCulloch, Minsky, Rochester et Holland, et d'autres. Des résultats partiels ont été obtenus mais le problème nécessite davantage de travail théorique.

4. Théorie de la taille d'un calcul

Si l'on nous donne un problème bien défini (un problème pour lequel il est possible de tester mécaniquement si une réponse proposée est valide ou non), une manière

de le résoudre est de tester toutes les réponses possibles dans l'ordre. Cette méthode est inefficace et pour l'exclure, on doit disposer d'un critère permettant de mesurer l'efficacité d'un calcul. Quelques réflexions montrent que pour obtenir une mesure de l'efficacité d'un calcul, il est nécessaire d'avoir en main une méthode pour mesurer la complexité des processeurs de calcul, ce qui peut être fait si l'on a une théorie de la complexité des fonctions. Des résultats partiels sur ce problème ont été obtenus par Shannon, et également par McCarthy.

5. Auto-amélioration

Probablement qu'une machine vraiment intelligente pourrait réaliser des tâches qui seraient mieux décrites en disant qu'elles permettent une auto-amélioration. Des schémas ont été proposés pour faire cela et méritent d'être mieux étudiés. Il semble que cette question puisse être étudiée de façon abstraite également.

6. Abstractions

Un certain nombre de types d'"abstraction" peuvent être définis distinctement et d'autres moins distinctement. Une tentative directe de les classer et de décrire les méthodes des machines pour fabriquer des abstractions à partir de perceptions et d'autres données sembleraient dignes d'intérêt.

7. Hasard et créativité

Une conjecture attractive et encore clairement incomplète est que la différence entre la pensée créatrice et la pensée compétente sans imagination réside dans l'injection d'un sain hasard. Le hasard doit être guidé par l'intuition pour être efficace. En d'autres termes, le fait de deviner d'une manière apprise ou bien l'intuition incluent l'un et l'autre une dose de hasard contrôlée dans une pensée ordonnée différemment.

Dans la proposition étaient fournies de courtes biographies des "proposants" :

Claude E. Shannon, Mathématicien, Laboratoires téléphoniques Bell. Shannon a développé la théorie statistique de l'information, l'application du calcul propositionnel aux circuits alternatifs, et a obtenu des résultats sur la synthèse efficace des circuits alternatifs, la conception de machines qui gagnent, la cryptographie, et la théorie des machines de Turing. Lui et J. McCarthy sont co-éditeurs des Annales de l'étude mathématique de la théorie des automates.

Marvin L. Minsky, professeur assistant en mathématiques et neurologie à Harvard. Minsky a construit une machine pour simuler l'apprentissage par réseaux de neurones et a soutenu une thèse en mathématiques à Princeton intitulée "les réseaux de neurones et le problème de la modélisation du cerveau" qui contient des résultats sur la théorie de l'apprentissage et sur la théorie des réseaux de neurones aléatoires.

Nathaniel Rochester, Directeur de recherches en théorie de l'information, IBM, à Poughkeepsie, New York. Rochester a étudié le développement de radars pendant sept ans et les machines informatiques pendant sept ans. Lui et un autre ingénieur ont été conjointement responsables de la conception du IBM type 701, qui est un ordinateur de grande dimension qui est largement utilisé aujourd'hui. Il a travaillé sur des techniques de programmation automatique largement utilisées aujourd'hui et a étudié comment l'on pourrait obtenir que des machines effectuent des tâches qui étaient jusque-là dévolues aux seuls humains. Il a aussi travaillé sur la simulation des réseaux de neurones avec un accent particulier sur l'utilisation des ordinateurs pour tester des théories en neurophysiologie.

John McCarthy, Professor assistant de Mathématique, Collège Dartmouth. McCarthy a travaillé sur un certain nombre de questions liées à la nature mathématique du processus de pensée incluant la théorie des machines de Turing, la vitesse des ordinateurs, la relation d'un modèle de cerveau à son environnement, et l'utilisation des langages par les machines. Quelques résultats de son travail sont inclus dans les Annales à venir, éditées par Shannon et McCarthy. Les autres travaux de McCarthy ont concerné le domaine des équations différentielles.

La proposition complète est disponible sur le site de Stanford. Les salaires des participants non supportés par des institutions privées (comme les laboratoires Bell et IBM) furent estimés à 1 200 \$ par personne, 700 \$ pour les deux étudiants diplômés. Les dépenses de train s'élevaient à 300 \$ et seraient remboursées aux participants qui ne résidaient pas dans le New Hampshire. La fondation Rockefeller prit en charge 7 500 \$ du budget estimé à 13 500 \$ pour la totalité de la conférence, qui dura pour tout dire six à huit semaines (les comptes-rendus diffèrent) à l'été 1956, commençant le 18 juin et se terminant le 17 août.

Participants

Bien qu'ils viennent d'une grande variété de domaines, que ce soient les mathématiques, la psychologie, l'ingénierie électrique et plus, les participants à la conférence de Dartmouth de 1956 partageaient une croyance qui les définissait en commun, notamment le fait que penser n'est pas une propriété spécifique que ce soit aux humains ou même aux êtres biologiques. Ils croyaient plutôt que le calcul est un phénomène de déduction formelle qui peut être compris de façon scientifique et que le meilleur instrument non humain pour effectuer une telle chose est l'ordinateur digital (McCormack, 1979).

Les quatre participants initiaux invitèrent chacun quelques personnes qui partageait leur croyance à propos de ces propriétés uniques de la cognition. Parmi elles il y avait les futurs lauréats du prix Nobel John F. Nash Jr (1928-2015) et Herbert A. Simon

(1916-2001). Le premier a vraisemblablement été invité par Minsky, et les deux ont été étudiants à Princeton ensemble vers le début des années 50. Simon a été invité par McCarthy lui-même ou par transitivité par Allen Newell, car tous deux avaient travaillé ensemble chez IBM en 1955. Selon les notes prises par Ray Solomonoff, vingt personnes ont participé sur la période des huit semaines. Parmi elles :

Herbert A. Simon (1916-2001)

Selon une des communications de McCarthy en mai 1956, Herbert A. Simon devait assister aux deux premières semaines du colloque.

Simon, alors Professeur de gestion à Carnegie Mellon (alors dénommée Carnegie Tech) avait jusque-là travaillé sur les problèmes de décision (qu'on appelait "comportement administratif") jusqu'à l'obtention de sa thèse en 1947 à l'Université de Chicago. Il irait jusqu'à gagner le prix Nobel d'Economie en 1978 pour son travail, et plus généralement pour "ses recherches pionnières dans les processus de décision mis en œuvre dans les organisations économiques". Au moment de la conférence, il collaborait avec Allen Newell et Cliff Shaw sur le langage-machine de traitement de l'information (IPL) et sur les programmes pionniers en théorie logique (Logic Theorist en 1956 et General Problem Solver en 1959). C'était le premier programme écrit pour simuler les capacités de résolution de problèmes d'un humain, et il réussissait à prouver 38 des 52 premiers théorèmes des Principia Mathematica de Whitehead et Russell, en trouvant même des preuves nouvelles et plus élégantes que celles proposées en 1912. Cliff travaillerait plus tard sur la machine universelle de résolution de problèmes, qui pourrait résoudre tout problème exprimé d'une manière suffisamment formelle par un ensemble de formules bien formées. Ce travail évoluerait en architecture Soar pour l'intelligence artificielle.

Allen Newell (1927-1992)

Allen Newell, le collaborateur de Simon, assista également aux deux premières semaines du colloque.

Ils s'étaient connus à RAND Corporation en 1952 et avaient créé ensemble le langage de programmation IPL en 1956. Ce langage était le premier à introduire la manipulation de listes, les listes de propriétés, les fonctions de haut niveau, le calcul symbolique et les machines virtuelles comprenant des langages de calcul digital. Newell était aussi le programmeur qui le premier introduisit le traitement des listes, l'application de l'analyse par les moyennes au raisonnement général et à utiliser des heuristiques pour limiter l'espace de recherche d'un programme. Un an avant la conférence, Newell avait publié un article intitulé "La machine qui joue aux échecs : un exemple de traitement d'une tâche complexe par adaptation", qui délimitait les contours théoriques d'

“..une conception imaginative pour qu'un ordinateur puisse jouer aux échecs d'une façon humanoïde, incorporant la notion de buts, les niveaux souhaités pour arrêter la recherche, se satisfaisant de mouvements “suffisamment bons”, des fonctions d'évaluation multi-dimensionnelles, la génération de

sous-buts pour implémenter la réalisation des buts et quelque-chose qui ressemblait à la “Best first search” (recherche du premier d’abord). L’information à propos de l’échiquier devait être exprimée symboliquement dans un langage qui ressemblait au calcul des prédicats.”

(citation extraite de “Allen Newell” dans les mémoires auto-biographiques d’Herbert A. Simon, 1997).

Newell attribue l’idée qui l’amena à son article à une “expérience de conversion” qu’il eut lors d’un séminaire en 1954 en écoutant un étudiant de Norbert Wiener, Oliver Selfridge des Laboratoires Lincoln Laboratories, décrire “un programme d’ordinateur qui apprend à reconnaître des lettres et d’autres formes” (Simon, 1997).

Oliver Selfridge (1926-2008)

Oliver Selfridge a rapporté avoir également assisté au colloque pendant deux semaines. Selfridge écrivait des papiers visionnaires importants sur les réseaux de neurones, la reconnaissance des formes, et l’apprentissage machine. Son “article Pandemonium” (1959) est considéré comme un classique dans le milieu de l’intelligence artificielle.

Julian Bigelow (1913-2003)

L’un des ingénieurs informaticiens pionniers était là lui-aussi. Bigelow avait travaillé avec Norbert Wiener sur l’un des papiers fondateurs au sujet de la cybernétique, et avait été recruté par John von Neumann pour construire l’un des premiers ordinateurs digital, à l’Institut des Etudes Avancées (IAS) en 1946 sur la recommandation de Wiener. Les autres participants incluaient Ray Solomonoff, Trenchard More, Nat Rochester, W. Ross Ashby, W.S. McCulloch, Abraham Robinson, David Sayre, Arthur Samuel et Kenneth R. Shoulders.



FIGURE 3 — Von Neumann et al.

Résultats

De façon assez ironique, les “échanges scientifiques intenses et soutenus pendant deux mois” imaginés avant la tenue du colloque par John McCarthy n’ont jamais vraiment eu lieu (McCormack, 1979) :

“La plupart de ceux qui étaient là étaient assez obnubilés par la poursuite des idées qu’ils avaient avant de venir, ou n’étaient pas là, et aussi loin que j’aie pu le voir, il n’y eut pas de réel échange d’idées. Les gens vinrent pour différentes périodes. L’idée était que chacun soit d’accord pour venir six semaines, et les personnes vinrent pour des périodes allant de deux jours à six semaines, ce qui fait qu’il n’y eut aucun moment où tout le monde était présent. C’était très décevant pour moi parce que cela signifiait vraiment que nous ne pourrions pas avoir de réunions régulières.”

Quelques résultats tangibles peuvent cependant être tirés de cette expérience. D’abord, le terme Intelligence artificielle (IA) lui-même fut inventé pour la première fois par McCarthy pendant la conférence. Parmi les travaux importants qui furent réalisés pendant la période de la conférence, McCarthy cita plus tard les travaux de Newell, Shaw et Simon sur le langage de traitement de l’information (IPL) et sur leur machine pour la théorie logique (Moor, 2006). Parmi d’autres résultats, un des participants à la conférence, Arthur Samuel, inventerait l’expression “apprentissage machine” en 1959 et créerait le programme de jeu d’échecs Samuel, l’un des premiers programmes auto-apprenant au monde. Oliver Selfridge est maintenant considéré comme le “père de la perception machine” pour ses recherches en reconnaissance des formes. Minsky serait récompensé par le prix Turing en 1969 pour son “rôle central dans la création, la mise en forme et les avancées apportées au domaine de l’intelligence artificielle”. Newell et Simon recevraient le prix Turing également en 1975 pour leurs contributions “à l’intelligence artificielle et à la psychologie cognitive humaine”.

Ci-dessous, un article résumant les événements de la conférence anniversaire AI@50 qui a été écrit après la conférence par James Moor et publié dans le Magazine AI Vol. 27 Numéro 4 en 2006.

La conférence de l’Intelligence artificielle du Collège de Dartmouth : les 50 années à venir

La conférence de l’Intelligence artificielle du Collège de Dartmouth : les 50 années à

venir (AI@50) a eu lieu du 13 au 15 juillet 2006.

La conférence avait trois objectifs :

- célébrer le projet de recherches d'été de Dartmouth, qui s'était tenu en 1956 ;
- dresser le bilan de la façon dont l'IA avait progressé ;
- et projeter la façon dont l'IA se développerait ou devrait le faire.

Cette conférence a été généreusement financée par le Doyen de la Faculté et le bureau du Prévôt du Collège de Dartmouth, par la Darpa, et par quelques donateurs privés.



FIGURE 4 — Cinq participants au projet initial devant la plaque de commémoration (lors du projet AI@50 en 2006) : Trenchard More, John McCarthy, Marvin Minsky, Oliver Selfridge, and Ray Solomonoff.

Réflexions sur 1956

Dater le début d'un mouvement est difficile, mais le projet de recherches d'été de 1956 est souvent considéré comme l'événement qui a fondé l'Intelligence Artificielle (IA) comme discipline de recherche. John McCarthy, alors professeur de mathématiques à Dartmouth, avait été déçu que les papiers sur les études concernant les automates, qu'il co-éditait avec Claude Shannon, n'en disent pas davantage sur les possibilités des ordinateurs de faire preuve d'intelligence. Aussi, dans la proposition écrite par John McCarthy, Marvin Minsky, Claude Shannon, et Nathaniel Rochester pour cet événement de 1956, McCarthy voulait, comme il l'a expliqué lors de la conférence AI@50, "clouer le drapeau au mât" (ne pas accepter la défaite). On attribue à McCarthy l'invention de l'expression "intelligence artificielle" ainsi que d'avoir solidifié

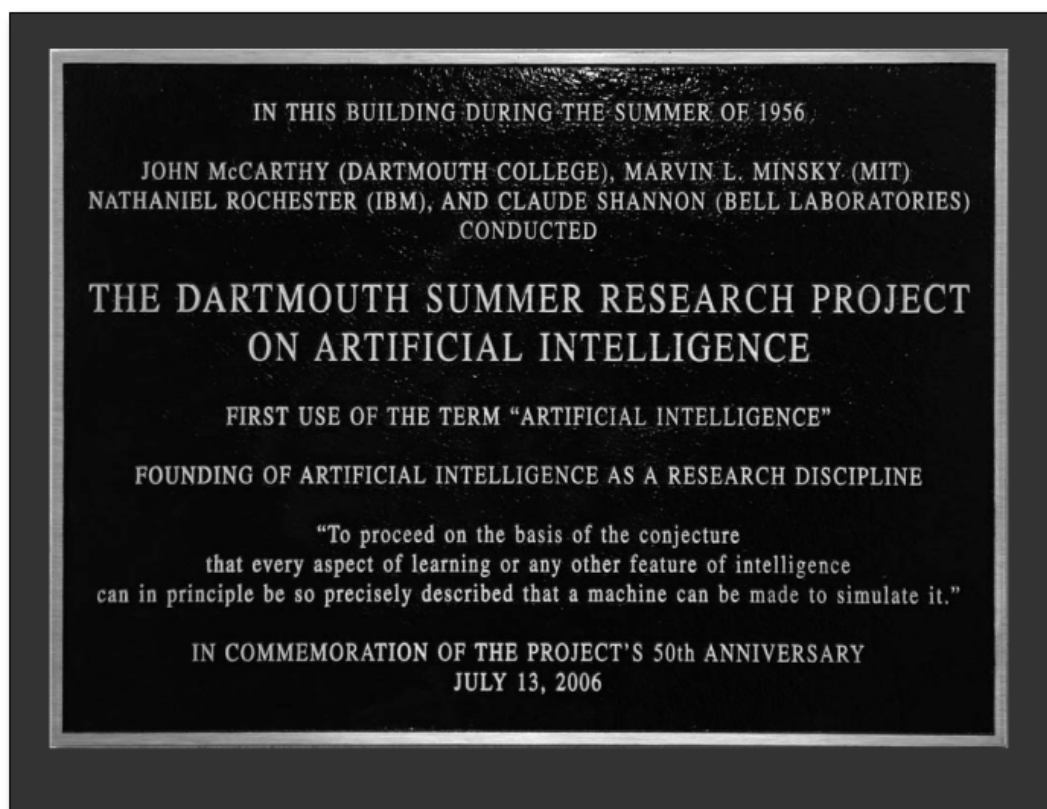


FIGURE 5 — Plaque commémorative.

profondément l'orientation de la discipline. Il est intéressant de spéculer pour savoir si le domaine aurait été différent si on l'avait appelé "intelligence calculatoire" ou avec un autre nom parmi de nombreux noms possibles.

Chacun donna ses souvenirs. McCarthy reconnut que le projet de 1956 ne donna pas naissance aux coopérations attendues. Les participants ne vinrent pas au même moment et la plupart restèrent fixés sur leur propre agenda de recherche. McCarthy insista sur le fait qu'il y avait néanmoins à ce moment-là d'importants développements de la recherche, particulièrement pour le langage IPL avec Allen Newel, Cliff Shaw et Herbert Simon sur la machine logique théorique.

Marvin Minsky fit comme commentaire que, bien qu'il ait travaillé sur les réseaux de neurones pour sa thèse quelques années avant le projet de 1956, il stoppa ce travail initial parce qu'il était convaincu que des avancées pourraient être faites avec d'autres approches informatiques. Minsky exprima ce constat que trop de personnes en IA aujourd'hui essayaient plutôt les approches populaires et ne publiaient que sur les succès.

Oliver Selfridge souligna l'importance des nombreux domaines de recherches avant et

après le projet d'été de 1956, ce qui a aidé l'IA en tant que domaine. Le développement des langages améliorés et des machines était essentiel. Il offrit des opportunités à de nombreuses activités pionnières telles que celles que J. C. R. Lickleiter développa plus tard, sur le temps partagé, ou bien la conception d'ordinateurs chez IBM par Nat Rochester ou encore le travail de Frank Rosenblatt sur les perceptrons.

Trenchard More fut envoyé sur le projet d'été pour deux semaines non consécutives par l'Université de Rochester. Quelques-unes des meilleures notes décrivant le projet IA ont été prises par More, bien qu'il admette ironiquement qu'il n'a jamais aimé l'utilisation des mots "artificielle" ou "intelligence" comme termes pour dénommer le domaine.

Ray Solomonoff dit qu'il vint sur le projet d'été pour convaincre tout le monde de l'importance de l'apprentissage machine. Il en repartit en ayant beaucoup perfectionné ses connaissances sur les machines de Turing, ce qui lui servit dans des travaux ultérieurs.

Ainsi, selon certains points de vue, le projet de recherche de 1956 ne répondit pas tout à fait à toutes les attentes. Les participants vinrent à différents moments, et travaillèrent sur leurs propres projets, et ainsi, ce ne fut pas tout à fait une conférence au sens habituel.

Il n'y eut pas d'accord sur une théorie générale du domaine, en particulier sur une théorie générale de l'apprentissage. Le domaine de l'IA a démarré non par accord sur la méthodologie ou par choix des problèmes ou de la théorie générale, mais par une vision partagée du fait que des ordinateurs pouvaient être fabriqués pour exécuter des tâches intelligentes. Cette vision a été hardiment établie dans la proposition de la conférence de 1956 : "L'étude doit être menée sur la base de la conjecture que tout aspect de l'apprentissage et de toute autre caractéristique de l'intelligence peut en principe être si précisément décrite qu'une machine doit pouvoir être capable de la simuler".

Evaluations en 2006

Il y eut plus de trois douzaines d'excellentes présentations et événements à la conférence AI@50, et il n'y a pas assez de place ici pour leur donner toute celle qu'elles méritent. Des chercheurs de haut niveau donnèrent des conférences au sujet de l'apprentissage, de la recherche, des réseaux, de la robotique, de la vision, du raisonnement, du langage, de la cognition et de la théorie des jeux.

Ces présentations montraient les résultats significatifs de l'IA durant le demi-siècle écoulé. Considérons la robotique comme un exemple. Comme l'a souligné Daniela

Rus, en 1956, il n'y avait pas de robots tels que nous les connaissons aujourd'hui. Il y avait des automates dédiés à des tâches spécifiques. Aujourd'hui, les robots sont partout. Ils nettoient nos maisons, explorent les océans, voyagent sur la surface de Mars, et gagnent le grand Challenge de la DARPA d'une course de 32 miles dans le désert de Mojave. Rus émet l'hypothèse que dans le futur, nous pourrions avoir nos propres robots personnels comme aujourd'hui nous avons nos ordinateurs personnels, des robots qui pourraient être programmés pour nous aider dans toute tâche que nous voudrions réaliser. On peut imaginer que des morceaux de robots puissent être assemblés pour devenir le genre de structure dont nous avons besoin à un moment donné. Beaucoup de progrès ont été accomplis en robotique, et beaucoup de progrès semblent atteignables dans un proche horizon.

Bien que l'IA ait obtenu de nombreux succès dans les 50 dernières années, de nombreux désaccords importants restent associés au domaine. Les différents domaines de recherche ne collaborent pas, les chercheurs utilisent des méthodologies de recherche différentes, et il n'y a toujours pas une théorie générale de l'intelligence et de l'apprentissage qui unifierait la discipline.

L'un des points de désaccord dont il a été débattu à AI@50 est la question de savoir si l'IA devrait être plutôt basée sur la logique, ou bien sur les probabilités.

McCarthy continue d'être un fervent partisan d'une approche basée sur la logique. Ronald Brachman a argumenté sur le fait qu'une idée maîtresse de la proposition pour le projet de 1956 était qu'"une grande partie des pensées humaines consiste à manipuler des mots selon certaines règles pour raisonner et certaines règles pour conjecturer" et que cette idée clef a servi de base commune à la plupart de l'IA des 50 dernières années. C'était la révolution de l'IA ou, comme l'a expliqué McCarthy, la contre-révolution, puisque c'était une attaque du comportementalisme, qui était devenu la position dominante en psychologie dans les années 1950.

David Mumford argumenta au contraire sur le fait que les 50 dernières années avaient expérimenté le déplacement depuis la fragile logique vers les méthodes probabilistes. Eugene Charniak alla dans le sens de ces arguments également en expliquant comment le traitement du langage naturel est maintenant un traitement essentiellement statistique. Il établit franchement : "Les statistiques ont pris le contrôle pour le traitement du langage naturel parce que ça marche."

Un autre point de désaccord, corrélé à la question *logique versus probabilités*, est le débat *psychologie contre paradigme pragmatique*. Pat Langley, dans l'esprit d'Allen Newell et Herbert Simon, a vigoureusement maintenu que l'IA devrait revenir à ses racines psychologiques si l'on souhaite obtenir une intelligence artificielle de niveau humain. D'autres chercheurs en IA sont plus enclins à explorer ce qui réussit même

fait par des moyens non humains. Peter Norvig suggéra que la recherche, particulièrement du fait de l'énorme quantité de données exploitables du web, pourrait montrer des signes encourageants de traitement des problèmes traditionnels d'IA même sans tenter de simuler une intelligence humaine. Par exemple, la traduction machine avec un degré raisonnable de précision entre l'arabe et l'anglais est maintenant possible en utilisant des méthodes statistiques alors même que personne dans le domaine de recherches en question ne parle telle ou telle langue que l'on a besoin de traduire.

Finalement, il y a eu le débat en cours sur l'utilité des réseaux de neurones pour atteindre l'intelligence artificielle. Simon Osindero qui travaille avec Geoffrey Hinton discuta des réseaux plus puissants. Terry Sejnowski et Rick Granger expliquèrent tous deux les nombreuses connaissances acquises sur le cerveau dans la dernière décennie et comment cette information était très pertinente pour construire des modèles en machine de l'activité intelligente.

Ces différences variées ont été utilisées pour montrer la bonne santé du domaine. Comme Nils Nilsson l'a souligné, il y a de nombreux chemins qui mènent au sommet. Bien sûr, toutes les méthodes peuvent ne pas être fructueuses sur la course longue, mais puisque nous ne savons pas quelle est la meilleure option, il est bon d'explorer différentes voies. Malgré les différences, comme en 1956, il y a une vision commune sur le fait que les ordinateurs peuvent exécuter des tâches intelligentes. Peut-être que cela seul unifie déjà la discipline.



FIGURE 6 — Collège de Dartmouth.

Prospective pour 2056

De nombreuses prédictions à propos du futur de l'IA ont été données à AI@50. Quand on leur a demandé ce que seraient les 50 prochaines années au regard d'aujourd'hui, les participants à la conférence initiale avaient différentes positions. McCarthy a fourni sa vision qui est qu'une IA de niveau humain n'est pas assurée pour 2056. Selfridge a exprimé que les ordinateurs feraient davantage que de la planification, incorporeraient des sentiments, et seraient affectés par eux, mais n'atteindraient pas le niveau humain. Minsky pensait que ce qui était nécessaire pour qu'il y ait un progrès significatif dans le futur était que quelques brillants chercheurs poursuivent leur recherche en poursuivant leurs propres bonnes idées, et non pas en faisant ce que leurs aînés avaient fait. Il s'est lamenté que trop peu d'étudiants aujourd'hui ne poursuivent de telles idées mais soient plutôt intéressés par l'entrepreneuriat ou la loi. La plupart espéraient que les machines seraient toujours sous le contrôle des humains et suggérèrent que les machines ne puissent jamais atteindre l'imagination des humains. Solomonoff prédit au contraire que des machines vraiment efficaces ne tarderont pas à être obtenues. Le danger, selon lui, est politique. Les technologies disruptives, comme le calcul, fournissent une grande puissance, puissance qui peut être mal utilisée, mise entre les mains d'individus et de gouvernements.

Ray Kurzweil offrit une vision bien plus optimiste à propos des progrès et déclara que nous devons avoir confiance dans le test de Turing, argument avec lequel beaucoup furent en désaccord. Prévoir les événements technologiques est toujours hasardeux. Simon avait prédit une machine championne d'échecs dans les 10 ans. Il a eu tort sur les 10 ans, mais cela arriva au bout de 40 ans. Ainsi, en se donnant une plus longue période, 50 ans de plus, il est fascinant de méditer à ce que l'IA pourra accomplir.

Sherry Turkle pointa judicieusement que la composante humaine est facilement négligée dans le développement technologique. Le résultat important pour nous concerne peut-être moins les capacités des ordinateurs que nos propres vulnérabilités quand nous serons confrontés à des intelligences artificielles très sophistiquées.

Plusieurs douzaines d'étudiants diplômés et post-doctoraux ont été financés par la DARPA pour assister à la conférence AI@50. Notre espoir est que beaucoup d'entre eux aient été inspirés parce ce qu'ils ont observé. Peut-être que certains d'entre eux présenteront leurs résultats à la conférence de célébration des 100 ans du projet de recherches d'été de Dartmouth.

Notes

Une plaque commémorant le projet de recherches de l'été 1956 a récemment été accrochée dans le Hall au Collège de Dartmouth, l'immeuble dans lequel les activités

ont eu lieu (figure 4).

Pour obtenir plus de détails sur les orateurs et les sujets, voir <https://www.dartmouth.edu/ai50/homepage.html>.

James Moor est professeur de philosophie au Dartmouth College. Il est professeur adjoint au Centre de philosophie appliquée et d'éthique publique (CAPPE) à l'Université Nationale Australienne. Il a obtenu sa thèse en Histoire et Philosophie des Sciences à l'Université d'Indiana. Ses domaines de recherche sont la Philosophie de l'intelligence artificielle, l'Éthique des machines, la Philosophie de la pensée, la Philosophie des sciences et la Logique. Il est éditeur du journal *Pensées et machines* et est président de la Société internationale de l'éthique de la technologie de l'information (INSEIT). Il a obtenu une récompense "Faire la différence" de l'association SIGCAS.