
**Programme de recherche
de l'ICSIA au CIFAR**

**Rétrospective de l'année 2025 :
bâtir une IA plus sûre au
service des Canadiennes
et Canadiens**

Table des matières

- 3 Introduction
- 4 Le programme de recherche de l'ICSIA au CIFAR en chiffres
- 6 Message de la direction
- 9 Protéger la Société
- 15 Rendre l'IA Fiable et Équitable
- 19 Assurer la Sécurité des Systèmes d'IA Critiques

[Le CIFAR](#) bénéficie du financement du gouvernement du Canada pour son rôle de direction de la Stratégie pancanadienne en matière d'IA, qu'il assume en collaboration avec l'[Amii](#), [Mila](#) et l'[Institut Vecteur](#).

Financé par le gouvernement du Canada

Canada

© 2026 CIFAR. Tous les droits sont réservés.

Introduction

En novembre 2024, le gouvernement fédéral a fondé l'Institut canadien pour la sécurité de l'IA (ICSIA), une initiative qui reflète le rôle crucial que joue la confiance du public dans la réussite de l'innovation. Relevant du ministère de l'Innovation, des Sciences et du Développement économique (ISDE) du Canada, l'ICSIA est une initiative fédérale dont l'organe scientifique indépendant est le Programme de recherche de l'ICSIA au CIFAR. Sa mission consiste à mobiliser les spécialistes en sécurité de l'IA de tous les coins du pays et de toutes les disciplines, afin de relever les défis d'ordre technique et sociétal soulevés par les systèmes d'IA avancés.

Le rapport Rétrospective de l'année 2025 : bâtir une IA plus sûre au service des Canadiennes et Canadiens dresse le bilan des progrès réalisés en 2025. En renforçant la capacité de recherche nationale et en favorisant l'établissement d'une masse critique de talents spécialisés, le CIFAR permet au Canada de se positionner en tant que chef de file mondial du développement de systèmes d'IA sûrs et dignes de confiance.

Le programme de recherche de l'ICSIA au CIFAR en chiffres

Au cours de sa première année d'existence, le programme de recherche de l'ICSIA au CIFAR a eu des retombées importantes sur la recherche en sécurité de l'IA, sur la formation et sur la mobilisation des connaissances et des idées grâce aux réalisations suivantes :

Création d'une communauté de recherche nationale sur la sécurité de l'IA

Financement d'un réseau de quelque 55 spécialistes de toutes les disciplines, dont :

28

chercheuses et chercheurs principaux et 6 membres du Conseil de recherches de l'ICSIA

11

postdoctorantes et postdoctorants en sécurité de l'IA au CIFAR

10

scientifiques, ingénieures et ingénieurs en sécurité de l'IA affiliés à l'Amii, à Mila ou à l'Institut Vecteur.

Réalisation de projets de recherche à fort impact

2,4 M\$ investis dans le lancement de
12 nouveaux projets de recherche, dont :

10

projets Catalyseur

2

Réseaux de solutions

Production de résultats

5

partenariats de recherche
d'envergure nationale et
internationale (en cours)

3

tables rondes sur la
sécurité de l'IA regroupant
responsables politiques
et spécialistes

28+

produits de
connaissances ou
résultats de recherche
attendus en 2026

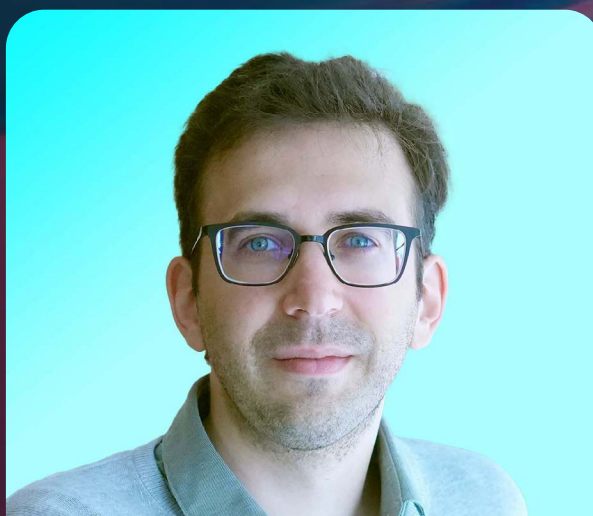
Message de la direction

L'année 2025 a marqué un tournant mondial important pour la sécurité de l'IA. À la suite des préoccupations urgentes soulevées par d'éminents spécialistes tels que Yoshua Bengio, titulaire d'une chaire en IA Canada-CIFAR, et Geoffrey Hinton, membre émérite du CIFAR — et à la suite de la parution du Rapport international sur la sécurité de l'IA —, le monde a reconnu la nécessité de concilier rapidité d'innovation et gestion rigoureuse des risques.



Catherine Régis

Codirectrice, Programme de
recherche de l'ICSIA au CIFAR //
Titulaire de chaire en IA Canada-
CIFAR, Mila, Université de Montréal



Nicolas Papernot

Codirecteur, Programme de
recherche de l'ICSIA au CIFAR //
Titulaire de chaire en IA Canada-
CIFAR, Institut Vecteur, Université
de Toronto

Le Canada était particulièrement bien préparé à relever ce défi. Chargé d'élaborer et de mettre en œuvre la toute première stratégie nationale sur l'IA au monde en 2017, le CIFAR a joué un rôle décisif dans la constitution d'un large vivier de talents et dans l'excellence en recherche qui fait aujourd'hui la réputation de notre pays. C'est sur ces bases solides qu'a été établi le Programme de recherche de l'ICSIA au CIFAR. Cette initiative nous permet de renforcer la capacité de recherche nationale et de favoriser l'établissement d'une masse critique de talents spécialisés pour permettre au Canada de se positionner en tant que chef de file mondial du développement de systèmes d'IA sûrs et dignes de confiance. Nous avons mis à profit ces acquis pour mobiliser rapidement le milieu scientifique canadien et passer de la théorie à des résultats tangibles en un an seulement.

En tant que codirectrice et codirecteur provenant de deux domaines distincts — le droit et l'informatique —, nous abordons la sécurité de l'IA comme un défi de nature sociotechnique. Notre programme fait le pont entre les dimensions techniques et sociales afin de garantir que les systèmes d'IA sont à la fois robustes et socialement responsables. Pour faire face à la profonde complexité des défis soulevés par la sécurité de l'IA, il est essentiel de soutenir des collaborations multidisciplinaires entre le milieu de la recherche, le gouvernement et d'autres parties prenantes.

Les résultats sont déjà au rendez-vous



Défense de la démocratie

élaborer des systèmes
propres à contrer
l'ingérence étrangère et
la désinformation



Protection de la jeunesse

concevoir des mécanismes
de détection et de
blocage des contenus
qui encouragent
les comportements
autodestructeurs



Sécurité du système judiciaire

lancer un Réseau de
solutions pour protéger les
tribunaux canadiens contre
les preuves synthétiques
produites par l'IA



📷 (De gauche à droite) Elissa Strome (CIFAR), Joel Martin (Conseil national de recherches Canada), Stephen Toope (CIFAR), Yoshua Bengio (Mila, LoiZéro), François-Philippe Champagne (Gouvernement du Canada), Valérie Pisano (Mila), Tony Gaffney (Institut Vecteur) et Cam Linke (Amii) lors de l'inauguration de l'Institut canadien de la sécurité de l'intelligence artificielle en novembre 2024.

Au cours de notre première année d'activités, nous avons financé 12 projets et soutenu 28 chercheurs et chercheuses de différentes disciplines. Notre programme fait avancer la recherche et l'innovation en permettant l'élaboration de nouvelles techniques et pratiques pour assurer la sûreté et la fiabilité de systèmes d'IA dont le déploiement respecte nos valeurs sociétales. En bref, il consiste à garantir que l'IA est digne de confiance.

Rien de cela ne serait possible sans le dévouement et l'engagement du milieu canadien de la recherche sur la sécurité de l'IA, du gouvernement fédéral et du ministère de l'Innovation, des Sciences et du Développement économique. Nous tenons à les remercier pour les efforts et les fonds qu'ils investissent sans relâche dans la recherche et les talents en sécurité de l'IA.

Nous remercions également les personnes et les organismes qui nous offrent leur soutien et leur collaboration, les instituts d'IA nationaux (l'Amii, Mila et l'Institut Vecteur), le Centre de recherches pour le développement international (CRDI), l'UK AI Security Institute et bien d'autres encore.

Ensemble, nous bâtissons une IA plus sûre au service des Canadiennes et Canadiens.

Protéger la Société

Pour protéger notre avenir collectif contre les risques d'envergure que posent les systèmes d'IA avancés, il faut atténuer des préjudices systémiques comme la désinformation massive et les bouleversements économiques. Il faut aussi mettre en place les outils et politiques nécessaires pour que l'IA demeure au service du bien commun.



Idées intelligentes avec Geoffrey Rockwell

Titulaire d'une chaire en IA Canada-CIFAR à l'Amii, Geoffrey Rockwell met à profit son expertise en éthique pour aborder la recherche sur la sécurité de l'IA d'un point de vue philosophique. Il explique le rôle que doit assumer le gouvernement pour réduire les risques et mobiliser les connaissances et infrastructures de sécurité existantes lors du déploiement de l'IA.

Pleins feux

Protéger la santé mentale contre les compagnons d'IA

De plus en plus de Canadiennes et Canadiens font appel à des robots conversationnels pour rompre l'isolement ou chercher une validation personnelle. Toutefois, les preuves s'accumulent concernant les effets néfastes sur la santé mentale d'un usage excessif ou inapproprié de ces outils. Les conséquences vont de la dépendance à la psychose induite par l'IA, en passant par l'incitation au suicide. [On estime que près de 70 %](#) des jeunes font régulièrement appel à des compagnons d'IA, d'où la nécessité de prévoir des mécanismes de protection indépendants tels que des filtres technologiques, des politiques d'utilisation et des programmes de sensibilisation.

La Programme de recherche de l'ICSIA au CIFAR finance le Studio de sécurité en intelligence artificielle de Mila afin de limiter les risques inhérents aux interactions nuisibles avec ces robots. Cette initiative vise à créer des mécanismes de protection fiables et indépendants et à concevoir des outils de référence complets, représentatifs de la diversité culturelle et sociale du Canada, afin de mesurer les préjudices en toute objectivité.

À ce jour, le Studio a développé une première version d'un système de protection de la santé mentale et d'un outil d'évaluation comparative des robots conversationnels propulsés par l'IA. Il poursuit ses travaux afin de les étendre à plusieurs fournisseurs de grands modèles de langage, ainsi qu'à plusieurs langues et particularités culturelles, en faisant appel à des données réelles anonymisées et à l'expertise de spécialistes en santé mentale.



« Le mécanisme de protection en santé mentale et les outils d'évaluation comparative de Mila permettront d'établir un moyen fiable et indépendant de mesurer l'ampleur des interactions nuisibles avec les compagnons d'IA, afin de protéger nos populations les plus vulnérables, dont nos enfants, contre l'incitation au suicide. »

Simona Grandabur

Responsable du Studio de sécurité en intelligence artificielle, Mila

Collaboration interdisciplinaire et regards vers l'avenir

« Ce qui m'emballe le plus dans ces travaux, c'est le soutien unanime d'un réseau de partenaires de différentes disciplines. Cette collaboration sociotechnique entre les membres des communautés touchées et des spécialistes de l'IA, de la santé mentale, des politiques publiques et de l'éducation assure l'établissement de mécanismes de protection solides et multidisciplinaires contre les effets néfastes des compagnons d'IA sur la santé mentale », déclare Simona Grandabur, responsable du Studio de sécurité en intelligence artificielle à Mila.

Au cours de la prochaine année, le Studio prévoit de développer des filtres intelligents pour bloquer les contenus générés par l'IA qui favorisent ou encouragent des comportements autodestructeurs. Il entend élaborer en outre des protocoles de tests de fiabilité pour évaluer la sécurité et la robustesse des modèles d'IA conversationnelle et générative. Le Studio concevra également des outils d'évaluation des risques psychologiques et éthiques.

La première version officielle du tableau de bord du Studio regroupant des outils d'évaluation comparative et des mécanismes de protection sera rendue publique en 2026.

Pleins feux

Protéger le Canada contre la désinformation

Les ingérences étrangères malveillantes et la désinformation alimentée par l'IA menacent directement la démocratie canadienne en cherchant à éroder la confiance envers nos institutions, nos médias et la société civile. En réponse à ce problème, un projet Catalyseur pour la sécurité de l'IA financé en 2025 par le CIFAR vise à développer un outil d'IA avancé destiné à protéger la population canadienne contre les campagnes de désinformation.

Ce projet de recherche visant à contrer les usages malveillants de l'IA est mené sous la houlette de Matthew E. Taylor, titulaire d'une chaire en IA Canada-CIFAR à l'Amii (Université de l'Alberta), de Brian McQuinn (Université de Regina) et de James Benoit, titulaire d'une bourse postdoctorale du CIFAR en sécurité de l'IA.

Un système de défense propulsé par l'IA

L'équipe travaille à mettre au point CIPHER, un système d'IA avancé qui maintient l'humain au centre de la prise de décision. Ce système vise à outiller les organisations de la société civile pour leur permettre de repérer et de déjouer des campagnes de désinformation coordonnées et complexes. Dans un premier temps, le projet se concentre sur la détection d'ingérences russes dans les contenus textuels et visuels afin de fournir un rempart essentiel à la société canadienne.

« Pour les membres du projet CIPHER, la fiabilité de l'information relève de la sécurité nationale. Repérer les campagnes de désinformation commanditées par des États permet à chacun et à chacune de ne pas perdre de vue les faits et les valeurs canadiennes. Nous voulons ainsi éviter que des influences extérieures viennent empoisonner nos débats et nos décisions de sécurité », a expliqué l'équipe au CIFAR.

Ce projet Catalyseur du CIFAR pour la sécurité de l'IA produira des résultats tangibles, dont :

- Une démonstration rigoureuse de la faisabilité de l'outil CIPHER, mis à l'essai dans des conditions réelles par des partenaires de la société civile au Canada et à l'international;
- Des documents d'orientation exploitables afin de guider les actions du gouvernement et de l'industrie;
- Un nouvel ensemble de données publiques visant à accélérer les nouvelles initiatives de recherche et développement dans ce domaine prioritaire afin d'étendre les retombées du projet au-delà de son champ d'application initial.



« L'IA est de plus en plus présente partout. Au lieu de lui confier les décisions importantes, notre architecture maintient systématiquement l'humain au cœur du processus. Le projet CIPHER vise à gagner la confiance des responsables politiques et du public afin qu'ensemble, nous puissions protéger les espaces démocratiques contre la désinformation et la mésinformation. »

Matthew E. Taylor

Titulaire de chaire en IA Canada-CIFAR, Amii



📷 (De gauche à droite) Kianna Adams (AltaML), Golnoosh Farnadi (Mila), Mohamed Abdalla (Amii) et Elissa Strome (CIFAR) discutent des moyens de concevoir des systèmes d'IA respectueux des valeurs et de la sécurité de la société ainsi que du rôle de chef de file que le Canada peut jouer à cet égard.

Projets financés

Réseau de solutions — Protéger les tribunaux contre les contenus synthétiques générés par l'IA

- Ebrahim Bagheri (Université de Toronto)
- Maura Grossman (Université de Waterloo)

Projet Catalyseur pour la sécurité de l'IA — Sur l'utilisation sécuritaire des modèles fondateurs à bruit statistique

- Mi Jung Park (titulaire de chaire en IA Canada-CIFAR, Amii, Université de la Colombie-Britannique)

Projet Catalyseur pour la sécurité de l'IA — CIPHER : Contrer l'influence par la mise en évidence de modèles et l'évolution des réponses

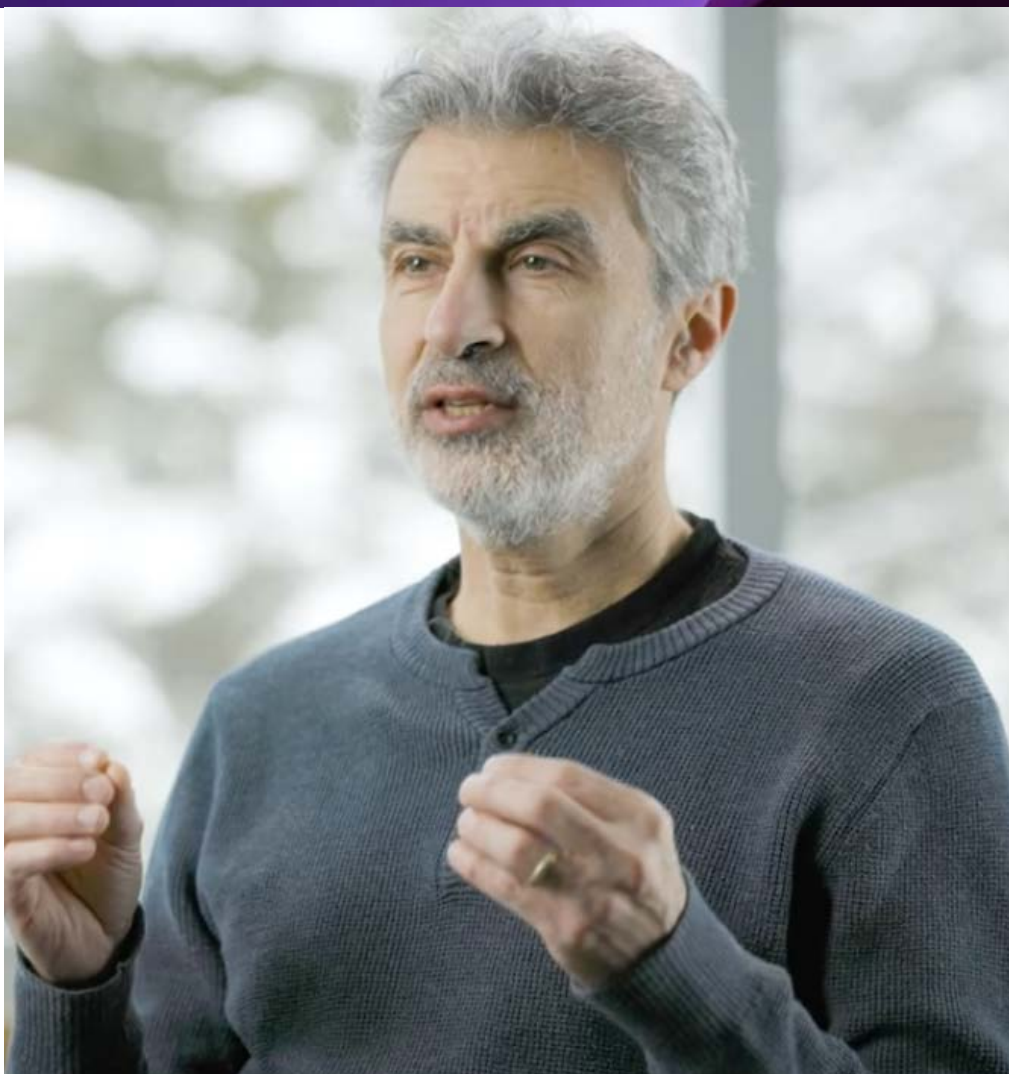
- Matthew E. Taylor (titulaire de chaire en IA Canada-CIFAR, Amii, Université de l'Alberta)
- Brian McQuinn (Université de Regina)
- James Benoit (titulaire d'une bourse postdoctorale du CIFAR sur la sécurité de l'IA, Amii)

Rendre l'IA Fiable et Équitable

Le fait d'intégrer activement les valeurs humaines et l'équité aux systèmes d'IA garantit que celle-ci reste au service du bien commun et qu'elle ne perpétue pas des préjudices sociaux comme les biais et la discrimination.

Idées intelligentes avec Yoshua Bengio

Comment concevoir des systèmes d'IA sans danger pour les gens? Yoshua Bengio, titulaire d'une chaire en IA Canada-CIFAR (Mila, Université de Montréal, LoiZéro), explique les différentes manières dont l'IA pourrait nuire à la société et pourquoi, malgré tout, il reste optimiste quant à l'avenir.



Pleins feux

Bâtir une IA plus sûre grâce à un raisonnement plus fiable

Comment faire confiance à une IA qui en sait plus que nous? Le volume de données traitées par les systèmes d'IA est tel que les utilisatrices et utilisateurs n'ont aucun moyen d'en vérifier les résultats. Pour remédier à ce problème, une solution envisagée réside dans le « débat entre modèles », un principe qui oblige les systèmes à citer leurs sources. Or, les modèles actuels ne peuvent pas argumenter de façon fiable : ils déforment les faits et sortent les sources de leur contexte. Pour être considérée comme fiable, l'IA doit apprendre à formuler et à justifier des raisonnements qui puissent être acceptés sans réserve.



« La forte interdisciplinarité de cette collaboration permet de repenser l'IA non seulement comme une question technique, mais aussi comme une question sociétale, politique et économique. Apprendre de spécialistes d'autres domaines autant que du sien est essentiel pour réaliser des progrès concrets dans la recherche actuelle en IA. »

Maria Ryskina

Titulaire d'une bourse postdoctorale du CIFAR sur la sécurité de l'IA

Dans le cadre de leur projet Catalyseur pour la sécurité de l'IA financé en 2025 par le CIFAR, Gillian Hadfield (Université Johns Hopkins, Université de Toronto [professeure auxiliaire], Institut Vecteur) et son équipe s'attaquent à cette question en mettant au point un cadre de débat en IA.

« Notre but est d'apprendre aux modèles de langage à formuler des justifications acceptables et à les étayer par des preuves valables tirées de sources d'information accessibles, explique Maria Ryskina, titulaire d'une bourse postdoctorale du CIFAR sur la sécurité de l'IA. Ces modèles inspireront davantage confiance, car leur supervision pourra être effectuée de manière fiable même par des profanes. »

La relation de confiance entre l'humain et l'IA

Tirant parti de connaissances issues du droit, de l'économie, de l'évolution culturelle et de la science politique, le cadre de débat est un projet visant à amener les systèmes d'IA à raisonner de manière normative, en prenant des décisions fondées sur les règles qui structurent les comportements — un peu comme les humains agissent selon les normes et les lois de la société. Ce projet vise à rendre les systèmes d'IA plus fiables, plus résilients et mieux arrimés aux structures sociales humaines.

L'équipe entend produire d'importants livrables techniques : nouvelles architectures d'agents d'IA, ensemble de données inédit portant sur un raisonnement normatif rigoureux, cadre de débat permettant de tester et d'appliquer des normes communes et invariables lors des interactions entre agents d'IA.

À terme, le but est d'améliorer la sécurité des systèmes d'IA en les dotant d'une meilleure connaissance des règles tacites et des concessions complexes qui régissent la société humaine dans toute sa diversité.

« La prolifération incontrôlée de l'intelligence artificielle a déjà des répercussions majeures dans la vie des gens. Je suis fier de prendre part à une initiative qui façonne l'avenir de l'IA de manière plus constructive — tant du point de vue du développement que de l'usage », ajoute Maria Ryskina.



 L'informaticienne de renom Deborah Raji (Mozilla) aborde les dimensions sociotechniques de la sécurité de l'IA à l'occasion de la première réunion annuelle du programme de recherche de l'ICSIA en octobre 2025.

Projets financés

Réseau de solutions — Réduire les biais dialectaux

- Laleh Seyyed-Kalantari (Université York)
- Blessing Ogbuokiri (Université Brock)

Projet Catalyseur pour la sécurité de l'IA — Échantillonnage d'explications latentes à partir de GML au profit d'un raisonnement sûr et interprétable

- Yoshua Bengio (titulaire de chaire en IA Canada-CIFAR, Mila, Université de Montréal, LoiZéro)

Projet Catalyseur pour la sécurité de l'IA — La robustesse antagoniste de la sécurité des GML

- Gauthier Gidel (titulaire de chaire en IA Canada-CIFAR, Mila, Université de Montréal)

Projet Catalyseur pour la sécurité de l'IA — Faire progresser l'alignement de l'IA par le débat et le raisonnement normatif partagé

- Gillian Hadfield (Université Johns Hopkins, Université de Toronto)
- Maria Ryskina (titulaire d'une bourse postdoctorale du CIFAR sur la sécurité de l'IA, Université de Toronto)

Projet Catalyseur pour la sécurité de l'IA — Formalisation des contraintes pour l'évaluation et l'atténuation du risque agentique

- Sheila McIlraith (titulaire de chaire en IA Canada-CIFAR, Institut Vecteur, Université de Toronto)

Assurer la Sécurité des Systèmes d'IA Critiques

Concevoir des outils rigoureux permettant d'évaluer la sécurité, l'exactitude et la fiabilité de l'IA avancée est indispensable pour protéger les systèmes critiques et favoriser l'innovation responsable au sein de l'industrie.

Idées de génie avec Nicolas Papernot

Comment bâtir des systèmes d'IA à la fois sûrs et dignes de confiance? Les travaux de Nicolas Papernot s'intéressent à la manière dont les machines peuvent assimiler des informations importantes sans compromettre les données personnelles des utilisatrices et utilisateurs. Protéger les systèmes critiques est essentiel pour gagner la confiance de la population canadienne, assurer la sécurité de leurs renseignements et accélérer l'adoption de l'IA.



Pleins feux

Une étude marquante évalue la sécurité et la fiabilité de l'IA de pointe

L'Institut Vecteur a réalisé une évaluation indépendante dans le cadre d'une étude comparative révolutionnaire visant à établir la sécurité et la fiabilité de grands modèles de langage (GML) utilisés partout dans le monde.

Piloté par l'équipe Ingénierie de l'IA de l'Institut Vecteur, le projet a évalué 11 modèles d'IA de pointe parmi les plus utilisés au monde. L'étude a porté sur des systèmes à code source ouvert et à code source fermé, dont la version de DeepSeek-R1 sortie en janvier 2025. Chacun a été soumis à un ensemble complet de 16 critères d'évaluation distincts.

Ce projet marque un jalon décisif dans la recherche sur la sécurité de l'IA. À mesure que la course mondiale à l'IA s'accélère, avec l'apparition de modèles de langage toujours plus puissants, il devient essentiel d'établir des outils comparatifs fiables et largement acceptés. Cette étude fournit un outil essentiel pour aider les scientifiques, les équipes de développement et les responsables politiques à évaluer la performance de ces modèles sur le plan de l'exactitude, de la fiabilité et de l'équité.

Favoriser l'adoption d'une IA fiable et sans danger

« Cette étude offre aux personnes et aux organisations un portrait clair et indépendant du comportement réel de ces modèles, afin que l'IA puisse être adoptée en toute sécurité et en toute confiance », affirme Deval Pandya, vice-président, Ingénierie de l'IA à l'Institut Vecteur.

Dans un souci de transparence et de reddition de comptes, l'Institut Vecteur a publié l'étude complète en libre accès, y compris ses résultats et le code sous-jacent. « En mettant nos travaux à la disposition du plus grand nombre, nous permettons à chacun et à chacune d'en valider les résultats, d'en tirer des apprentissages et de les bonifier en vue d'une utilisation plus intelligente et plus sûre de l'IA partout au pays », ajoute Deval Pandya.



« Je me réjouis de contribuer à un avenir où l'IA inspirera confiance aux gens parce qu'ils pourront en voir et en tester le fonctionnement par eux-mêmes. En publiant nos évaluations et nos outils en libre accès, nous permettons à une vaste communauté d'en reproduire les résultats, de repérer les lacunes et de les corriger. C'est la clé d'une adoption sécuritaire, non seulement en laboratoire, mais aussi dans les hôpitaux, les salles de classe, les entreprises et les services publics. »

Deval Pandya

Vice-président, Ingénierie de l'IA, Institut Vecteur

L'évaluation repose sur de puissants outils comparatifs comme MMLU-Pro, MMMU et OSWorld, aujourd'hui largement utilisés dans le domaine. Conçus par Wenhua Chen et Victor Zhong, professeurs à l'Université de Waterloo et titulaires de chaires en IA Canada-CIFAR dont [les travaux](#) enrichissent notre approche des techniques d'évaluation comparative, ces outils sont maintenant largement adoptés par des géants de l'industrie tels qu'OpenAI, Google et Anthropic.



📷 Blessing Ogbuokiri (Université Brock), coresponsable du Réseau de solutions Réduire les biais dialectaux, lors de la première réunion annuelle du programme de recherche de l'ICSIA en octobre 2025.

Projets financés

Projet Catalyseur pour la sécurité de l'IA — Des laboratoires de chimie autonomes et sécuritaires

- Alán Aspuru-Guzik (titulaire de chaire en IA Canada-CIFAR, Institut Vecteur, Université de Toronto)

Projet Catalyseur pour la sécurité de l'IA — La robustesse antagoniste dans les graphes de connaissances

- Ebrahim Bagheri (Université de Toronto)
- Jian Tang (titulaire de chaire en IA Canada-CIFAR, Mila, HEC Montréal et Université de Montréal)
- Benjamin Fung (Mila, Université McGill)

Projet Catalyseur pour la sécurité de l'IA — Maintenir un contrôle significatif : comment concilier agentivité et surveillance dans le codage assisté par l'IA?

- Jackie Chi Kit Cheung (titulaire de chaire en IA Canada-CIFAR, Mila, Université McGill)
- Jin Guo (Université McGill)

Projet Catalyseur pour la sécurité de l'IA — Assurance de la sécurité et ingénierie pour les systèmes d'IA multimodaux reposant sur des modèles fondateurs

- Foutse Khomh (titulaire de chaire en IA Canada-CIFAR, Mila, Polytechnique Montréal)
- Lei Ma (titulaire de chaire en IA Canada-CIFAR, Amii, Université de l'Alberta)
- Randy Goebel (Amii, Université de l'Alberta)



CIFAR

LÀ OÙ L'AVENIR PREND SON **ÉLAN**
CIFAR.CA/FR