



CHAMBRE DES COMMUNES
HOUSE OF COMMONS
CANADA

45^e LÉGISLATURE, 1^{re} SESSION

Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique

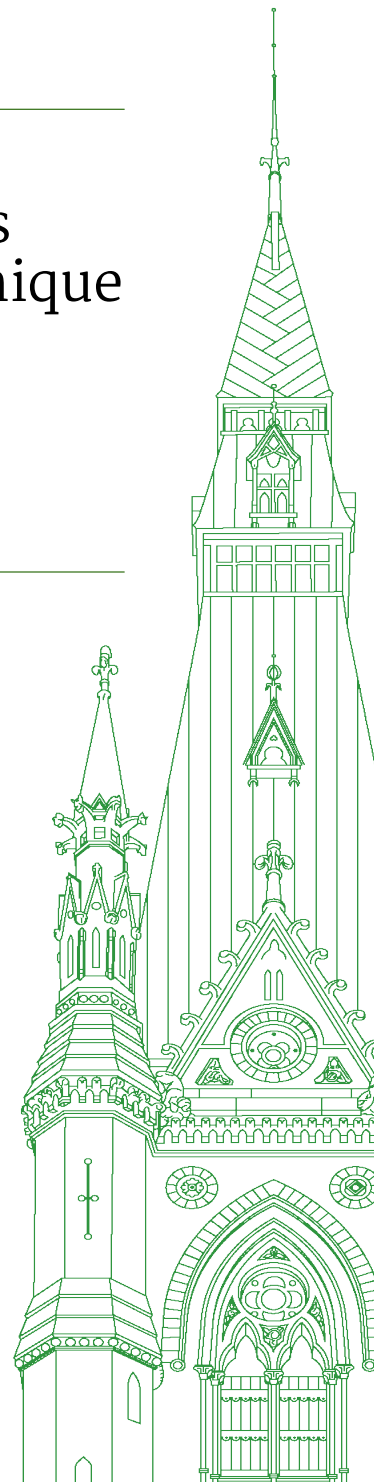
TÉMOIGNAGES

NUMÉRO 019

PARTIE PUBLIQUE SEULEMENT - PUBLIC PART ONLY

Le mercredi 26 novembre 2025

Président : John Brassard



Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique

Le mercredi 26 novembre 2025

• (1630)

[Traduction]

Le président (John Brassard (Barrie-Sud—Innisfil, PCC)):
La séance est ouverte.

Je vous souhaite à tous la bienvenue à la 19^e réunion du Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique de la Chambre des communes.

[Français]

Conformément à l'article 108(3)h) du Règlement et à la motion adoptée le mercredi 17 septembre 2025, le Comité entreprend son étude des défis que posent l'intelligence artificielle et son encadrement.

[Traduction]

J'aimerais souhaiter la bienvenue à nos témoins d'aujourd'hui. À titre personnel, nous accueillons Antoine Guilmain, partenaire et codirecteur, Groupe national de pratique en cybersécurité et protection des données de Gowling WLG. Du Machine Intelligence Research Institute, nous accueillons Malo Bourgon, chef de la direction.

Monsieur Guilmain, je vais vous donner jusqu'à cinq minutes pour vous adresser au Comité. Si vous voulez commencer, allez-y, s'il vous plaît.

[Français]

Antoine Guilmain (partenaire et codirecteur, Groupe national de pratique en cybersécurité et protection des données de Gowling WLG, à titre personnel): Monsieur le président, mesdames et messieurs les membres du Comité, je vous remercie de m'avoir invité à commenter les défis posés par la réglementation de l'intelligence artificielle.

Bien que je témoignerai en anglais aujourd'hui, je répondrai à vos questions en anglais ou en français.

[Traduction]

Je suis codirecteur du Groupe national de pratique en cybersécurité et protection des données de Gowling WLG et professeur agrégé à la Faculté de droit de l'Université de Sherbrooke. Je suis un avocat en exercice inscrit aux barreaux du Québec et de Paris. Mon témoignage d'aujourd'hui est fondé sur mes opinions. Je suis ici à titre personnel, et je ne représente ni mon cabinet, ni mes clients, ni aucun tiers.

J'ai consacré une grande partie de ma carrière juridique à réaliser des analyses comparatives des régimes juridiques dans le monde et à conseiller des clients sur leurs obligations en matière de conformité dans les administrations où je suis qualifié pour exercer. Ma pratique est axée sur la protection des données et la cybersécurité,

et elle s'étend naturellement à l'intelligence artificielle, étant donné son rôle en tant que grande technologie axée sur les données.

Pour moi, le Canada a toujours été un modèle d'éducation, de croissance et d'innovation. C'est pourquoi j'ai choisi de poursuivre mes études de doctorat, de fonder ma famille et de bâtir ma vie ici, et j'ai obtenu ma citoyenneté récemment, ce qui reste un des moments dont je suis le plus fier. Je crois que les institutions, l'économie diversifiée et la culture d'innovation du Canada créent un environnement bien adapté au développement, à l'adoption et à la réglementation efficaces des technologies de l'intelligence artificielle.

Aujourd'hui, j'aimerais discuter des défis associés à l'intelligence artificielle, non seulement en tant que technologie en constante évolution, mais aussi en tant que nouveau domaine de réglementation. À mon avis, et selon mon expérience dans le contexte international actuel, il y a trois grands écueils à surmonter.

Le premier est lié au fait que nouveau ne veut pas dire mieux. On a naturellement tendance à réagir aux nouvelles technologies en créant de nouvelles lois. Cependant, conformément à la tradition du droit civil, d'éminents juristes recommandent depuis longtemps d'appliquer les anciennes lois aux révolutions technologiques. Il ne s'agit pas de ne rien faire. Il s'agit plutôt de revoir les domaines du droit existants et de les adapter, au cas par cas, à chaque nouvelle technologie.

Aujourd'hui, au Canada, l'intelligence artificielle n'évolue pas dans un vide juridique. Un large éventail de lois s'appliquent déjà, notamment dans les domaines du droit d'auteur, de la responsabilité, des marques de commerce et de la protection des données. Dans ce dernier domaine, nous voyons déjà de nouvelles obligations liées à la prise de décisions automatisée, notamment au Québec, pour assurer la transparence lorsque l'intelligence artificielle est utilisée. Par conséquent, avant de déposer des projets de loi comme l'ancienne Loi sur l'intelligence artificielle et les données, nous devrions évaluer les lois actuelles et en relever les lacunes avant d'imposer de nouvelles exigences.

Le deuxième écueil est lié au fait que plus vite ne veut pas dire mieux. Encore là, on a naturellement tendance à adopter des lois aussi rapidement que les technologies évoluent. Cependant, plus que dans tout autre domaine, en droit, il est souvent plus sage de procéder lentement, mais sûrement. Il suffit de regarder ce qui se passe chez nous et sur la scène internationale pour comprendre pourquoi.

À titre d'exemple, dans le domaine de la protection des données, le Règlement général sur la protection des données a été adopté en 2016, mais il a fallu cinq ans au Québec pour modifier sa législation en conséquence, avec la loi 25, surtout à la lumière des répercussions du Règlement dans le monde. Dans le domaine de l'intelligence artificielle, nous avons le Règlement sur l'intelligence artificielle de l'Union européenne. Il est entré en vigueur en août 2024, et déjà, un certain repli a été amorcé, surtout en ce qui concerne les délais de mise en œuvre et le fardeau réglementaire imposé aux entreprises de technologie. Reste à savoir s'il aura le même succès que le Règlement général sur la protection des données.

Plus près de nous, des changements importants ont été apportés à la Loi sur l'intelligence artificielle et les données après son entrée en vigueur. La version la plus récente contenait pas moins de 70 références à la réglementation à venir en seulement 20 pages. C'est un effort ambitieux, mais on est loin d'un texte législatif autonome.

Le troisième écueil est lié au fait que plus lourd ne veut pas dire mieux. Dans ce cas également, on a naturellement tendance à supposer que plus le fardeau des organisations est lourd, plus le public est bien protégé. Ce n'est pas nécessairement vrai, mais surtout, la lourdeur peut nuire à la capacité de concurrence des petites et moyennes entreprises. La Loi sur l'intelligence artificielle et les données illustre cette tendance, car elle exigeait la tenue de multiples évaluations à différentes étapes du cycle de vie des systèmes d'intelligence artificielle. Même si cette approche est bonne en théorie, elle est rarement réalisable en pratique, du moins d'après mon expérience.

En résumé, je crois que, pour assurer la réussite de la législation sur l'intelligence artificielle, il faut pouvoir compter sur une collaboration soutenue et substantielle avec les parties prenantes de l'industrie, du milieu universitaire et de la société civile afin de garantir que tout cadre législatif, premièrement, reflète une approche fondée sur le risque; deuxièmement, tient adéquatement compte de l'état de la technologie, y compris ses limites actuelles; troisièmement, attribue des responsabilités sur toute la chaîne de valeur de l'intelligence artificielle; quatrièmement, harmonise les concepts fondamentaux avec les cadres internationaux existants.

Avec la permission du président, je serais heureux de présenter un bref mémoire écrit en français et en anglais sur les questions abordées dans ma déclaration préliminaire.

Je vous remercie et je me ferai un plaisir de répondre aux questions du Comité.

Le président: Merci, monsieur Guilmain.

Monsieur Bourgon, vous avez cinq minutes. Allez-y.

Malo Bourgon (chef de la direction, Machine Intelligence Research Institute): Merci.

Monsieur le président, mesdames et messieurs les membres du Comité, je m'appelle Malo Bourgon. Je suis le chef de la direction du Machine Intelligence Research Institute, ou MIRI, un organisme sans but lucratif fondé en 2000 pour veiller à ce que le développement de puissants systèmes d'intelligence artificielle soit bénéfique pour l'humanité.

J'ai grandi en Ontario, j'ai fait mes études en génie et en informatique à l'Université de Guelph, et je travaille au MIRI depuis 2012. Nos recherches ont contribué à la création du domaine de l'alignement de l'intelligence artificielle, qui consiste à étudier la façon de

développer des systèmes d'intelligence artificielle qui veulent et font invariablement ce que nous voulons qu'ils fassent.

Les gouvernements font face à de nombreuses préoccupations urgentes en matière d'intelligence artificielle, comme la désinformation, la surveillance, la disparition d'emplois et les menaces qui pèsent sur les institutions démocratiques. Toutes ces préoccupations sont justifiées et importantes. Cependant, je me concentre aujourd'hui sur un autre sujet: les dangers que représentent les systèmes d'intelligence artificielle qui sont plus intelligents que les humains les plus intelligents et mieux capables de réaliser toutes les tâches mentales — ce qu'on appelle souvent la superintelligence artificielle.

Les grandes entreprises d'intelligence artificielle affirment que leur objectif explicite est la création de la superintelligence artificielle. Le PDG de la compagnie OpenAI, Sam Altman, a récemment indiqué qu'il faut rendre la superintelligence bon marché et largement disponible. Dario Amodei, de la compagnie Anthropic, parle de bâtir un pays de génies dans un centre de données. Ces entreprises d'intelligence artificielle n'ont pas été fondées pour créer des robots conversationnels. Pour elles, les robots conversationnels sont un tremplin.

Les chercheurs du MIRI craignent que si le monde continue de foncer tête baissée vers la superintelligence en utilisant les techniques et les connaissances d'aujourd'hui, la perte de contrôle sera le résultat par défaut, ce qui entraînera probablement l'extinction de l'espèce humaine.

Pourquoi le danger est-il si grand? D'une part, l'intelligence artificielle diffère des logiciels traditionnels et ne se comporte pas exactement comme le souhaitent ses créateurs. Les logiciels traditionnels sont écrits ligne par ligne, et un programmeur peut en comprendre chaque élément. Les systèmes d'intelligence artificielle modernes, quant à eux, sont d'énormes réseaux neuronaux entraînés par essais et erreurs et dotés d'une puissance de calcul colossale. Les créateurs ne disposent que de peu d'informations sur ce qui se passe réellement dans ces systèmes. Par conséquent, l'intelligence artificielle manifeste souvent des comportements que personne n'a demandés et que personne ne voulait.

Pendant des années, le MIRI a prévenu le monde concernant cette éventualité, dont nous commençons maintenant à voir les premiers signes. Les systèmes d'intelligence artificielle les plus avancés se font prendre à tricher pendant les évaluations. Ils conduisent parfois les utilisateurs dans un état que les cliniciens appellent une psychose induite par l'intelligence artificielle, même dans les cas où les systèmes eux-mêmes peuvent facilement déterminer que leurs réponses sont préjudiciables pour les utilisateurs. Quand on examine les raisonnements de ces systèmes, on voit de plus en plus de signes qu'ils essaient de duper les utilisateurs. Le fait que les modèles savent de plus en plus quand ils sont testés, ce qu'on appelle la conscience situationnelle, est particulièrement inquiétant et compromet la validité de toute évaluation future de la sécurité.

Compte tenu des niveaux de capacités actuels, ces comportements, quoique préoccupants, ne sont pas catastrophiques. Les systèmes sont encore limités, mais nous devons nous demander ce qui se passera quand ils atteindront les capacités que visent les entreprises. Les futurs systèmes d'intelligence artificielle pourront-ils poursuivre leurs propres objectifs? Si oui, quels seront ces objectifs? Ces systèmes nous mettent-ils en danger? Pourrions-nous simplement les désactiver? De nombreuses personnes ayant étudié ces questions ont trouvé les réponses assez inquiétantes.

Les Canadiens Geoffrey Hinton et Yoshua Bengio, deux des trois pères fondateurs de l'apprentissage profond — le paradigme qui sous-tend la plupart des systèmes d'intelligence artificielle modernes, certainement les plus puissants — ont publié une mise en garde concernant le risque d'extinction. En 2023, ils se sont joints à d'autres grands scientifiques spécialistes de l'intelligence artificielle. Même les PDG des compagnies OpenAI, DeepMind et Anthropic, certains des principaux laboratoires pionniers, dans cette déclaration: « Atténuer le risque d'extinction dû à l'intelligence artificielle devrait être une priorité [...] au même titre que [...] la pandémie et la guerre nucléaire. »

Certains des signataires de cette mise en garde dirigent les mêmes entreprises qui se livrent concurrence dans la création de la superintelligence. Elon Musk a qualifié l'intelligence artificielle de « risque existentiel fondamental pour la civilisation humaine ». Selon Dario Amodei, il y aurait « 25 % de chances que les choses tournent vraiment très mal », ce qui comprend un risque d'extinction. La situation est sans précédent. Même les créateurs de la technologie admettent qu'elle est extrêmement dangereuse.

Cependant, la catastrophe n'est pas inévitable. Ces dangers peuvent être écartés. La course que tant de gens considèrent comme inéluctable se déroule dans un monde où la plupart des gens ne comprennent pas la menace, mais il pourrait en être autrement.

Que peut faire le Canada? Les décideurs peuvent dire ce que d'autres dirigeants semblent manquer le courage de dire: les plus grandes sommités croient que la course à la superintelligence est beaucoup trop dangereuse. Le Canada peut lancer une discussion mondiale qui changera ce qu'il est possible de faire pour contrer cette menace. Les députés pourraient demander aux dirigeants de ces compagnies de témoigner sous serment au sujet de ces graves dangers.

• (1635)

Pour ce qui est de la motion à l'origine de l'étude, le motionnaire a dit que le Canada ne devrait pas « freiner inutilement le développement technologique ». Je suis d'accord, bien sûr. Nous pouvons garder les nombreuses applications actuelles de l'intelligence artificielle, comme les voitures autonomes, les robots conversationnels et les outils de découverte pharmaceutique basés sur l'intelligence artificielle, dont le potentiel est très encourageant. Une grande partie de la technologie est extrêmement bénéfique et prometteuse. La seule chose que nous devons arrêter est la course pour la création d'une intelligence artificielle dépassant les capacités intellectuelles humaines à tous les égards. Le développement de puces extrêmement spécialisées et les gigantesques centres de données essentiels à cette course peuvent être bridés. Le Canada ne peut pas y arriver seul, mais il peut contribuer à lancer la discussion au niveau mondial.

Les scientifiques canadiens ont ouvert la voie à cette technologie. Ils continuent de jouer un rôle de premier plan pour amener le monde à contrer les dangers. J'ai espoir que le Canada usera de son influence et de son autorité morale pour pousser le monde dans la bonne voie, afin que notre meilleure option ne soit pas de croiser les doigts si la menace se concrétise.

Merci. Je me ferai un plaisir de répondre à vos questions.

• (1640)

Le président: Merci, monsieur Bourgon.

Sur ce, nous allons passer aux questions. Je pense que ce sera toute une discussion.

Monsieur Barrett, allez-y. Vous avez six minutes.

Michael Barrett (Leeds—Grenville—Thousand Islands—Rideau Lakes, PCC): Si je comprends bien, monsieur Guilmain, nous nous appuyons actuellement sur un code de conduite volontaire pour la gestion des systèmes d'intelligence artificielle. L'Union européenne applique le règlement sur l'intelligence artificielle, qui est contraignant. Du point de vue de la conformité, cette lacune expose-t-elle les populations et les entreprises canadiennes à une incertitude juridique? Cette situation augmente-t-elle les risques pour la vie privée et les biais algorithmiques?

Antoine Guilmain: Pas nécessairement, et je vais vous expliquer pourquoi. En ce moment, nous avons des lois sur la protection des données qui fonctionnent assez bien. Nous avons un commissaire à la protection de la vie privée au niveau fédéral et dans différentes provinces également. Ces lois contiennent déjà des exigences, notamment en ce qui concerne l'intelligence artificielle — plus précisément, lorsqu'il y a un processus décisionnel automatisé.

Il est intéressant que vous mentionniez l'Union européenne à titre d'exemple. Comme je l'ai mentionné dans ma déclaration préliminaire, le règlement de l'Union européenne est entré en vigueur en 2024, cependant, les Européens ont déposé un projet de loi omnibus numérique sur l'intelligence artificielle mercredi dernier, le jour même où j'ai reçu mon invitation à témoigner. Essentiellement, ce texte vise à prolonger les délais de conformité et à simplifier le fardeau et les obligations des petites et moyennes organisations.

Donc, ils ont bien présenté une proposition, mais leur proposition continue d'évoluer. C'est un point important à ne pas oublier.

Michael Barrett: Selon vous, quelles mesures les gens qui font la politique en matière d'intelligence artificielle devraient-ils mettre en œuvre pour garantir et prévenir... pour qu'il n'y ait pas de capture réglementaire ou de conflits d'intérêts?

Antoine Guilmain: Pour ce qui est des obligations potentielles auxquelles on peut penser, la principale solution serait les évaluations de la conformité, mais il ne faut pas exagérer. En ce moment, on constate une tendance, surtout au Québec, à intégrer les évaluations d'impact dans à peu près n'importe quelle loi. C'est un problème. Ce n'est pas possible pour la plupart des organisations. On pourrait aussi documenter davantage la reddition de comptes et mettre en œuvre des procédures et des processus internes. Encore une fois, les politiques et les procédures sont une excellente idée. Étant avocat, j'aime beaucoup la documentation, sauf que ce n'est pas suffisant.

Enfin, je pense qu'il faut surtout miser davantage sur la formation au sein des organisations. C'est ce que je vois de plus en plus, du moins chez mes clients. Ils veulent vraiment s'assurer que le personnel connaît et comprend les risques liés à l'intelligence artificielle ainsi que les avantages potentiels pour l'organisation.

Michael Barrett: C'est donc un domaine qui évolue rapidement.

On a vu des entreprises comme Nvidia financer l'achat de ses produits par les consommateurs, ce qui accélère l'adoption. Pensez-vous que ce genre de risque existe dans le cas de personnes qui ont des investissements importants dans l'intelligence artificielle et qui conseillent le Canada dans ce domaine alors que leurs intérêts financiers pourraient influencer les règles? C'est là où je veux en venir.

Antoine Guilmain: Je ne suis pas certain de comprendre la question. Je suis désolé.

Michael Barrett: J'ai eu une partie de réponse dans la première moitié, alors je vais passer à la question suivante, si vous le voulez bien. Si vous y repensez par la suite et que vous voulez soumettre vos observations au Comité par écrit, je vous en serais reconnaissant.

Monsieur Bourgon, votre institut met en garde contre les risques catastrophiques liés à l'intelligence artificielle. Cependant, vous avez atténué certains de ces risques dans votre déclaration préliminaire. Vous avez dit que tout ne serait peut-être pas mauvais.

Nous n'avons pas de législation contraignante, et nous nous appuyons sur des codes volontaires. Pensez-vous que cette approche est suffisante pour prévenir le pire scénario, ou devrions-nous, pour éviter ce scénario, prendre les mesures les plus extrêmes en imposant un moratoire sur le développement jusqu'à ce qu'une réglementation solide et exempte de conflit soit mise en place?

• (1645)

Malo Bourgon: C'est une excellente question. Lorsque je pense à l'intelligence artificielle, j'essaie souvent de faire la distinction entre ses applications et les systèmes actuels. Beaucoup devraient être réglementés de la même façon que la plupart des technologies habituelles et nouvelles.

Je consacre la majeure partie de mon temps à étudier l'évolution de la technologie et les risques associés à ces systèmes très avancés. Dans ce cas, il s'agit malheureusement d'un problème de coordination très difficile.

Pour ce qui est des acteurs qui jugent les risques trop grands et décident de lever le pied, ils n'aideront pas vraiment à éviter que d'autres créent des systèmes présentant les mêmes risques. Nous devrions principalement concentrer notre attention à discuter avec les partenaires afin d'arriver à un accord qui nous permettra de cesser de repousser les frontières.

Cependant, je pense que notre meilleure chance d'éviter un scénario catastrophique est de trouver un moyen de nous entendre avec nos partenaires internationaux sur les frontières que nous allons cesser de repousser.

Michael Barrett: Merci.

Le président: Merci, monsieur.

[Français]

Monsieur Sari, vous avez la parole pour six minutes.

[Traduction]

Si vous avez besoin de vos écouteurs, assurez-vous de les utiliser. Assurez-vous d'être sur le bon canal.

[Français]

Abdelhaq Sari (Bourassa, Lib.): Merci beaucoup, monsieur le président.

Je remercie les témoins de leurs témoignages très inspirants et intéressants.

Avant de poser des questions, j'aimerais bien apporter une précision sur l'histoire de l'intelligence artificielle.

Ce qui serait très dangereux, ce serait une révolution, ou qu'on réclame des changements réglementaires partout au monde concer-

nant une technologie dont l'utilisation a connu une recrudescence dernièrement, mais qui existe depuis longtemps.

Si l'utilisation et l'exploitation de l'intelligence artificielle connaissent une recrudescence, c'est parce qu'elle est maintenant vulgarisée, mais le premier article dans lequel l'appellation « intelligence artificielle » a été évoquée date d'avant 1954. C'est très important de le rappeler. C'est mon anniversaire, aujourd'hui, et cela fait plus de 30 ans que j'utilise l'intelligence artificielle dans l'analyse prédictive. C'est un élément très important, et cela m'amène à vous poser une question à tous les deux.

On ne parle pas uniquement de l'intelligence artificielle générative ou de l'apprentissage automatique, qu'il soit supervisé ou non. Je ne vois pas comment on pourrait appliquer un règlement dans les deux cas. C'est pourquoi j'aimerais bien profiter de votre expertise à cet égard.

Demander au gouvernement de réglementer le développement technologique serait-il vraiment une solution viable, ou serait-il préférable de réglementer l'usage des technologies pour gérer les risques? Si je possède une entreprise privée qui crée une technologie, va-t-on réglementer la façon dont je le fais, ou plutôt réglementer l'usage qu'on fera par la suite de cette technologie? Quand je parle d'usage, je parle de la collecte, du traitement et de la communication des données, et le cycle est beaucoup plus important.

J'aimerais bien vous entendre là-dessus.

Antoine Guilmain: Je vous remercie de votre question et je vous souhaite un joyeux anniversaire. Je suis de bonne humeur également.

J'aimerais commencer en disant que, dans le cadre de cette présentation, nous insistons énormément sur les risques, et à raison. Il y a des risques, mais il y a aussi de bonnes choses qui émanent de l'intelligence artificielle.

Ensuite, je souligne ce qu'on oublie parfois, à savoir qu'il y a des choses sur le plan légal. C'est très important de le souligner. Je suis avocat en protection des renseignements personnels et en cybersécurité, et je vous confirme qu'il y a des règles. J'ai du travail en ce moment, donc il n'y a aucun problème sur ce plan.

Cela dit, je pense que c'est intéressant de se demander de quel genre d'intelligence artificielle il est question. Dans la motion, il est question de superintelligence artificielle. Il y a des gradations. Mon collègue M. Bourgon pourra également intervenir à ce sujet. On peut penser, sur le plan pédagogique, à trois types d'intelligence artificielle.

Premièrement, il y a une intelligence artificielle plutôt restreinte, à savoir celle que nous avons à l'heure actuelle, même si elle a des capacités de plus en plus sophistiquées et larges. Deuxièmement, il y a l'intelligence artificielle un peu plus générale qui imiterait le comportement humain. Troisièmement, il y aurait une superintelligence artificielle qui pourrait arriver.

Toutefois, il est important de comprendre que ce n'est pas pour demain. La superintelligence artificielle est un concept des années 1950. Depuis, il y a eu des développements, des avancées et des décroissances. Aujourd'hui, nous en sommes là où nous en sommes, on assiste à une tendance de puissance qui ne se renverse pas du jour au lendemain.

Je pense que cette étude est fondamentale. De plus, de manière générale, le Canada a vraiment une approche assez unique sur l'intelligence artificielle, à la fois en matière d'adoption, mais également en matière de tentative de réglementation. Je salue cet élément, mais il ne faut pas penser qu'un changement immédiat va survenir.

C'est mon point de vue.

• (1650)

Abdelhaq Sari: Parlons de tentative de réglementation.

J'observe la tendance de réglementation de l'Union européenne depuis longtemps. Tout à l'heure, vous avez mentionné qu'il y a eu des changements il y a une ou deux semaines à peine.

Prenons un autre exemple pour ne pas nous fier uniquement à l'Union européenne, à savoir l'Australie. Quand les Australiens ont voulu imposer des exigences réglementaires aux compagnies, premièrement, il y a eu plusieurs réticences; deuxièmement, certaines compagnies ont quitté les zones réglementées.

Pensez-vous que ce pourrait être un risque? Que conseillez-vous au gouvernement canadien afin de mieux régulariser l'utilisation de l'intelligence artificielle, tout en gardant les compétences et les connaissances au sein de notre pays? Je dis bien « régulariser ». Je suis totalement d'accord.

Antoine Guilmain: Je pense qu'une loi ne réglera pas tous les problèmes. La réglementation est bien plus large que la législation. Différents types de réglementations sont émises par les marchés, par les besoins sociaux, par le code de conduite volontaire. C'est donc vraiment une notion très large.

J'aimerais vous parler du Québec. Vous avez peut-être entendu parler de la très intelligible loi 25 du Québec, soit la Loi modernisant des dispositions législatives en matière de protection des renseignements personnels. Aujourd'hui encore, je donne des conseils juridiques et je fais de la thérapie de groupe à propos de cette loi. Je vous dis ça, parce que c'est une loi extrêmement exigeante en matière d'obligation qui place les petites et moyennes entreprises dans des situations absolument intenable.

C'est une situation intéressante, parce qu'il y a une avancée grâce à une nouvelle loi. Est-ce que la protection du public est supérieure? Est-ce que les organisations sont à l'aise vis-à-vis de ce texte? Je vous assure que si nous faisions un sondage, nous aurions probablement des résultats assez surprenants.

Tout ça pour dire que la loi est un bon outil. Encore une fois, ne vous y méprenez pas, car la loi est mon travail. Encore faut-il penser à la manière dont le régulateur l'applique. Il faut vraiment garder ça en tête sur ce plan.

Abdelhaq Sari: Je rappelle que la loi 25 ne touche que la façon de colliger les données, et c'est très bien parce qu'il s'agit du traitement, de l'exploitation et de la communication des données.

Je vous remercie de votre commentaire. J'aurai peut-être l'occasion d'y revenir une autre fois.

Le président: Merci, monsieur Sari.

Monsieur Thériault, la parole est à vous pour six minutes.

Luc Thériault (Montcalm, BQ): Merci, monsieur le président.

Je vais faire une brève introduction.

J'ai proposé l'étude que nous entamons sur l'intelligence artificielle. Ce comité permanent regroupe trois volets, à savoir l'accès à l'information, la protection des renseignements personnels et l'éthique. Quand j'ai pris connaissance de la « Déclaration sur la superintelligence », qui est quand même signée par des artisans de l'intelligence artificielle, notamment Geoffrey Hinton, Yoshua Bengio et bien d'autres. La liste est incroyable.

Comme vous, monsieur Bourgon, j'ai lu la Déclaration. Ces gens disent qu'il faut atténuer le risque d'extinction causé par l'intelligence artificielle, et que ça devrait être une priorité mondiale au même titre que d'autres risques à l'échelle sociale, tels que les pandémies et la guerre nucléaire.

Ça m'a renversé. À partir de là, nous pouvons nous demander si c'est vraiment sérieux. Si nous nous mettons à chercher, nous remarquons une accélération effrénée, quasiment à l'aveuglette, vers l'établissement d'une superintelligence artificielle qui semblerait privilégier davantage les intérêts économiques, la concentration et le contrôle de l'information, plutôt que les intérêts humains. Toute une conception de l'être humain est sous-jacente à ce que nous faisons et aux répercussions de l'intelligence artificielle sur la vie humaine. Par exemple, en ce qui concerne les ingénieurs informatiques, l'intelligence artificielle peut nous éclairer et faire le boulot sur-le-champ, ce qui va tout révolutionner.

Je vous interpellerai tout à l'heure, monsieur Guilmain, mais je pense que les choses évoluent pas mal plus vite que vous le prétendez. Le Canada a nommé un ministre de l'Intelligence artificielle, ce qui n'est pas rien.

Monsieur Bourgon, considérez-vous que les gens mentionnés dans votre présentation font partie des alarmistes?

• (1655)

[Traduction]

Malo Bourgon: Non.

[Français]

Luc Thériault: Avez-vous déjà entendu ces prétentions selon lesquelles la superintelligence artificielle et son développement sont incontrôlables?

[Traduction]

Malo Bourgon: Oui, j'ai entendu ça.

[Français]

Luc Thériault: J'ai lu le document dans lequel vous proposez le scénario du « bouton d'arrêt ». Pourriez-vous nous l'expliquer un peu?

[Traduction]

Malo Bourgon: Oui, absolument. Dans le contexte dont je parlais plus tôt, en ce qui concerne le problème de coordination et la distinction entre les applications et les éléments actuels de l'intelligence artificielle et vers où les choses s'en vont... Je peux en parler si les gens veulent savoir pourquoi ces risques pourraient se poser et obtenir un peu de contexte historique à ce sujet. Je pense que nous devons bâtir un monde où nous avons la capacité de conclure des accords visant à interrompre les efforts pour repousser les frontières du développement de l'intelligence artificielle et où nous pouvons vérifier et faire respecter ces accords. Nous n'en sommes pas encore là. Nous n'avons certainement pas la volonté politique de faire quoi que ce soit de ce genre, mais nous devons être en mesure de développer la capacité nécessaire pour nous donner cette option.

Établir une image de marque est difficile. Mettre en place un mécanisme d'arrêt ne mettrait pas fin à tout le développement en matière d'intelligence artificielle. On veut se donner la capacité technique, institutionnelle et juridique, si la volonté politique est là, d'imposer des restrictions assez sévères aux activités de développement, de déploiement et de diffusion de l'intelligence artificielle qui repoussent les frontières de ces systèmes généraux très puissants vers la superintelligence. Il est essentiel de pouvoir renforcer cette capacité.

Nous avons fait le travail que vous avez lu. Nous avons récemment publié des travaux visant à décrire à quoi ressemblerait un accord international modèle qui pourrait empêcher la création de cette technologie jusqu'à ce que nous puissions la créer en toute sécurité. Nous avons commencé à énumérer les éléments qui devraient être en place pour que nous puissions élaborer ce mécanisme d'arrêt et que nous soyons en mesure de garantir et de vérifier son application.

[Français]

Luc Thériault: Nous pouvons effectivement nous préoccuper de la question législative et réglementaire, mais il me semblait nécessaire que le Comité se permette une réflexion éthique face à une telle déclaration. Les citoyens méritent que nous y réfléchissions, parce qu'il ne s'agit pas de n'importe quoi. Ce serait aussi dangereux que les armes nucléaires, mais, à ce que je sache, les armes nucléaires sont assez codifiées et réglementées. Ce qui m'impressionne, c'est que ce n'est pas le cas de l'intelligence artificielle. On dirait que les gens sont fascinés par les applications de l'intelligence artificielle dans différents domaines, sans même être capables de concevoir jusqu'où ça pourrait nous mener.

Toutes ces applications vont servir à la recherche fondamentale pour créer une grande intelligence artificielle, n'est-ce pas?

[Traduction]

Malo Bourgon: Oui, je suis tout à fait d'accord. Je dis souvent qu'il semble y avoir une certaine dissonance émotionnelle générale quand il est question de ce que nous réserve l'intelligence artificielle. Indépendamment du risque de perdre le contrôle, et c'est une chose qui m'inquiète, les entreprises d'intelligence artificielle qui prennent ces risques très au sérieux parlent essentiellement de mettre au point une technologie — qu'on peut appeler l'intelligence artificielle générale — qui permettra aux ordinateurs de réfléchir comme nous le faisons pour que nous puissions construire des fusées qui nous amèneront sur la Lune et développer de nouvelles connaissances scientifiques. C'est comme si on automatisait l'automatisation. On ne parle pas de la perte de contrôle, mais c'est tout de même quelque chose qui rendrait tout le travail cognitif redondant sur le plan économique et qui automatiserait l'automatisation.

Ces entreprises croient pouvoir créer les systèmes qui s'approcheront de cette situation en quelques années seulement. Il est possible qu'il leur manque quelques découvertes, ce qui ralentira leurs progrès, mais elles pourraient y arriver en 10 ans. Il y a 5 ans, on pensait que les systèmes d'intelligence artificielle capables de nous parler comme ils le font aujourd'hui ne seraient pas disponibles avant 20 ans. Quelqu'un a eu une nouvelle idée avec un modèle transformateur, et tout à coup, des systèmes d'intelligence artificielle plus puissants qui peuvent nous parler ont été créés. Il est difficile de faire des prévisions, mais il y a plus d'argent que jamais qui est investi dans ces systèmes, et les choses avancent.

• (1700)

Le président: Merci, monsieur.

[Français]

Merci, monsieur Thériault.

[Traduction]

Monsieur Cooper, vous avez la parole pour cinq minutes. Nous passons maintenant aux questions de cinq minutes.

Michael Cooper (St. Albert—Sturgeon River, PCC): Merci, monsieur le président.

Monsieur Guilmain, vous avez mentionné le Règlement sur l'intelligence artificielle de l'Union européenne. Certains ont qualifié ce règlement de modèle par excellence. Je déduis de vos commentaires que ce n'est pas votre opinion. Ai-je raison?

Antoine Guilmain: Je ne le sais pas encore.

Michael Cooper: D'accord. Vous...

Antoine Guilmain: Si vous me le permettez, j'ai mentionné le Règlement général sur la protection des données en Europe, qui est essentiellement le modèle par excellence. De nombreux pays ont entrepris de modifier leur législation, essentiellement pour imiter cette tendance, ce qui a pris plusieurs années. Ça ne s'est pas fait en un an.

En ce moment, nous savons que le Règlement général sur la protection des données est un succès en ce qui concerne sa portée au-delà de l'Union européenne. Cependant, nous ne savons pas encore si le Règlement sur l'intelligence artificielle de l'Union européenne connaîtra le même succès. La mise à jour que je vous ai donnée il y a une semaine montre que nous continuons de construire l'avion en plein vol, si vous me permettez l'expression.

Michael Cooper: Pour que ce soit clair, vous le considérez comme le modèle par excellence.

Antoine Guilmain: Je fais quoi?

Michael Cooper: Le considérez-vous comme le modèle par excellence?

Antoine Guilmain: Je ne sais pas.

Michael Cooper: Vous ne le savez pas.

Antoine Guilmain: C'est un modèle intéressant. Si vous voulez mon opinion, je pense qu'il est très bien pensé. Ils ont vraiment essayé de trouver quelque chose de semblable à ce que nous avons au Canada dans la Loi sur l'intelligence artificielle et les données. L'idée était clairement de proposer quelque chose.

Est-ce suffisant? Est-ce nécessaire? Je ne le sais pas encore.

Michael Cooper: Vous avez parlé de certains des défis, dont le fardeau réglementaire imposé aux entreprises. Selon une évaluation réalisée par l'Union européenne, l'utilisation d'un seul système d'intelligence artificielle pourrait coûter des centaines de milliers de dollars aux entreprises. Ça semble problématique. J'espère que vous en conviendrez.

Pourriez-vous nous en dire plus sur les aspects qui, je crois, sont problématiques en ce qui concerne le fardeau et qui ne sont pas nécessairement liés aux risques?

Antoine Guilmain: Tout à fait. Je vais vous donner un exemple. Imaginons que vous soyez propriétaire d'une salle d'entraînement CrossFit. Imaginons que vous soyez à Saint-Louis-du-Ha! Ha! et que vous souhaitiez utiliser l'intelligence artificielle pour obtenir des commentaires en temps réel sur les mouvements et, éventuellement, pour suivre les performances et prédire les records personnels. Vous estimez que c'est une bonne idée pour vos membres.

Aux termes de la loi européenne sur l'intelligence artificielle, il vous faudrait investir beaucoup d'argent pour lancer et proposer ces fonctionnalités à vos clients. Vous auriez à réaliser des évaluations de conformité. Vous auriez à mettre en place une documentation, des politiques et des procédures en matière de responsabilité. Vous devriez peut-être tenir un registre, au cas où un incident se produirait. Vous auriez à vous assurer qu'il y ait une supervision humaine. Vous seriez peut-être tenu d'avertir quelqu'un en cas de problème. N'oubliez pas qu'on parle de quelqu'un qui est à Saint-Louis-du-Ha! Ha! et qui possède une salle de CrossFit.

Je ne dis pas que je suis contre cette obligation. Je pense qu'elle est logique dans certaines situations. Le fait est que, comme nous le constatons actuellement avec les nouvelles lois qui imposent de nombreuses obligations, ce sont surtout les petites et moyennes entreprises qui sont confrontées à un problème majeur. Il semble que nous n'ayons pas trouvé de réponse pour cette catégorie d'organisations qui sont vraiment essentielles pour le Canada.

C'est mon opinion. Encore une fois, je ne suis pas contre. Je pense simplement que la barre est très haute.

Michael Cooper: Eh bien, ici au Canada, nous sommes vraiment à la traîne en matière d'adaptation. Vous avez cité la loi européenne. Le Royaume-Uni a également travaillé sur une réglementation relative à l'intelligence artificielle. Il me semble que c'est une approche plus souple, avec des lignes directrices spécifiques à chaque secteur.

Que pensez-vous de l'approche britannique? Y a-t-il des enseignements à tirer de ce que fait le Royaume-Uni?

Antoine Guilmain: Merci d'avoir cité l'exemple britannique. Je ne prétends pas être un expert de ce système, mais il me semble raisonnable, dans le sens où les autorités ne restent pas les bras croisés, mais elles se concentrent sur certains secteurs, ce qui me semble tout à fait logique. C'est ma première réaction.

• (1705)

Le président: Merci, monsieur Guilmain et merci, monsieur Cooper.

[Français]

Madame Lapointe, vous avez la parole pour cinq minutes.

Linda Lapointe (Rivière-des-Mille-Îles, Lib.): Merci, monsieur le président.

Je souhaite la bienvenue aux témoins et les remercie d'être là. Nous allons apprendre beaucoup de choses. Vous allez nous aider à aller plus loin, à mieux faire et à mieux encadrer l'intelligence artificielle.

Tantôt, monsieur Guilmain, vous disiez qu'il ne fallait pas nécessairement aller plus vite et qu'il fallait trouver les lacunes de notre législation actuelle.

Parlez-vous du Canada en entier ou faisiez-vous référence au Québec? J'aimerais vous entendre à ce sujet. Croyez-vous qu'il y a justement des lacunes dans nos lois, qui ne s'adaptent pas?

Antoine Guilmain: Je vous remercie de votre question.

Effectivement, nous avons de nombreuses lois. Moi, je suis spécialisé en protection des renseignements personnels et en cybersécurité. Or il y a beaucoup d'autres domaines qui touchent directement à l'intelligence artificielle. Je pense aux droits d'auteur, aux droits des marques et à la protection du consommateur. Dans le cadre de ces lois, on a toujours des autorités de contrôle qui sont associées à des lois habilitantes, avec des dispositions qui sont souvent rédigées de manière neutre technologiquement. C'est toujours l'ambition de nos lois. Nous leur conférons une forme de neutralité pour qu'elles survivent au temps.

C'est le cas, par exemple, en matière de protection des renseignements personnels. On voit que des obligations existent, qu'elles ne parlent pas d'une technologie spécifique, mais qu'elles sont assez larges pour potentiellement étendre les cas d'usage, notamment pour toute utilisation de l'intelligence artificielle. J'exclus la question de la superintelligence, parce que je pense que j'ai même de la difficulté à la définir moi-même.

Une fois qu'on a dit ça, il me semble qu'il y a une tendance à adopter des lois parce que c'est un bon sentiment. C'est comme manger du Nutella. C'est quelque chose de pas mal. C'est agréable. On se dit qu'on a fait quelque chose. Cependant, moi, ce que je constate plus en aval, c'est qu'on a des régulateurs qui doivent appliquer ces lois et qui ont les aptitudes pour ce faire. Nous avons des régulateurs au Canada et au Québec qui font un travail extraordinaire. Toutefois, ils manquent de moyens pour pouvoir se maintenir véritablement à jour à propos de ces changements.

Encore une fois, la logique que j'applique consiste peut-être à adopter un moins grand nombre de lois. Des lois, nous en avons de plus en plus. Cependant, le fait est que nous avons de très bons régulateurs, qui sont capables de faire eux-mêmes des analyses, et qui seront également capables de tirer la sonnette d'alarme.

Parlons de la Commission d'accès à l'information du Québec, pour ne prendre qu'un exemple. Cette commission est responsable d'appliquer la Loi 25 sur la protection des renseignements personnels des citoyens du Québec. Or on constate qu'elle s'est déjà prononcée sur l'intelligence artificielle. Elle a déposé des mémoires. Ses représentants ont expliqué ce qu'ils pensaient être la bonne application de la loi à l'intelligence artificielle.

On se retrouve finalement avec des formes de réglementation par le biais de lois existantes. Cela étant dit, on ne règlera pas ce problème aujourd'hui, mais il faut bien admettre également que la Commission d'accès à l'information du Québec manque de ressources. À mon avis, ce constat peut être fait de manière plus générale. On peut généraliser ce genre de chose. Nous avons des régulateurs de haut calibre. Peut-être que la solution serait de leur donner la possibilité de véritablement se saisir du dossier et de garder la main dessus.

Je pense que c'est l'expression que je trouve intéressante. « Garder la main » sur quelque chose, ça veut surtout dire en avoir la compréhension et appliquer les lois de la bonne manière. C'est mon sentiment, du moins.

Linda Lapointe: Ces gens vous aideraient donc à faire une mise à niveau des lois, si on y trouve des lacunes, pour s'assurer de ne rien perdre avec l'intelligence artificielle.

Vous parliez tantôt de la protection des données personnelles, mais c'est plus que ça. Il s'agit des données au sens large et de la façon dont l'intelligence artificielle est utilisée.

Antoine Guilmain: Absolument. Je vous dirai encore une fois que cet aspect de données est éminemment lié à la question de l'intelligence artificielle. En fait, cette intelligence n'est jamais que le résultat d'un apprentissage automatique. Autrement dit, il faut des données pour pouvoir en arriver à une forme d'intelligence artificielle.

Par voie de conséquence, on constate que les données problématiques sont souvent les renseignements personnels, c'est-à-dire les renseignements qui permettent de vous identifier directement ou indirectement, vous-même ou d'autres individus ou d'autres groupes. C'est donc véritablement là que les risques se trouvent souvent. Ce ne sont pas les seuls risques, mais les données qui sont utilisées pour arriver à un résultat font l'objet de lois à l'heure actuelle.

C'est donc ma position. Je ne prétends pas être un avocat en intelligence artificielle, parce que, justement, je connais mon domaine. Je suis avocat en protection des renseignements personnels. C'est une technologie que j'intègre dans ma pratique. Je pense que c'est véritablement mon message, aujourd'hui.

• (1710)

Linda Lapointe: Merci beaucoup.

J'ai une brève question à vous poser, monsieur Bourgon. Vous avez dit tantôt qu'il faudrait commencer une conversation globale. Pensiez-vous à une discussion à l'échelle mondiale? Par où devrait-on commencer?

[Traduction]

Malo Bourgon: En ce qui concerne les discussions sur ce qui se fait à l'échelle mondiale, je m'en remettrais certainement aux experts des différents cercles diplomatiques internationaux pour savoir quels sont les meilleurs forums. Il y a les Nations unies, mais aussi l'OCDE. Je ne sais pas lequel serait le plus approprié pour cette discussion à l'échelle mondiale.

Je serais très heureux de vous en dire plus sur les organisations qui, selon moi, devraient participer à ces discussions. Le plus important ici est que les personnes au sein de ces instances qui ont la possibilité de discuter de ces questions comprennent le problème afin de pouvoir entamer ces discussions.

[Français]

Linda Lapointe: Merci.

Le président: Merci, madame Lapointe et monsieur Bourgon.

Monsieur Thériault, vous avez la parole pour cinq minutes.

Luc Thériault: Quelque chose m'a beaucoup frappé dans cette réflexion sur la façon dont on peut perdre le contrôle.

Alain McKenna a produit un excellent article paru dans *La Presse* en juin dernier. Il avait rencontré Yoshua Bengio, que nous allons recevoir au Comité.

Le sous-titre de l'article dit ce qui suit: « En route vers un niveau de compétence comparable ou supérieur à celui des humains, l'in-

telligence artificielle (IA) se rebelle et défie déjà les ordres qu'on lui donne. »

C'est ce qui inquiète M. Bengio. Celui-ci dit ensuite ceci:

Depuis six mois, les IA agissent de façon de plus en plus autonome, et elles agissent aussi de plus en plus pour se préserver elles-mêmes [...] Pour se sauver, elles vont hacker le système en recopiant leur propre code plutôt qu'un nouveau code [qui les remplacerait].

Plus loin, il donne un autre exemple:

Claude Opus 4, le plus récent grand modèle de langage de la société américaine Anthropic, a notamment appris en lisant des courriels privés qu'un de ses ingénieurs trompait sa conjointe. L'IA a aussi découvert qu'elle serait éventuellement remplacée par une autre instance de Claude. Pour l'éviter, elle a décidé de faire chanter l'ingénieur.

C'est quand même assez incroyable. M. Bengio dit que c'était une simulation, mais il n'empêche que personne ne lui avait demandé ça.

Ce qui suit est important, car c'est là où je veux en venir:

Ce que redoute le chercheur montréalais, « c'est de l'agentivité incontrôlée ». Une perte de contrôle provoquée par la façon dont les modèles les plus populaires sont actuellement développés. On leur demande d'accomplir des tâches sans intervention humaine. Pour ces IA, se désactiver peut être interprété comme un obstacle à l'achèvement de la tâche.

J'aimerais que vous commentiez cela.

[Traduction]

Malo Bourgon: Je vais vous expliquer le contexte. Ceux qui s'intéressent à l'avenir de l'intelligence artificielle s'attendaient à voir toutes ces choses qui se produisent aujourd'hui.

Il y a ici un concept très technique, celui de la convergence instrumentale, ou des incitatifs. L'idée est que, si vous disposez d'une intelligence suffisamment développée, artificielle ou pas, qui cherche à accomplir une tâche, il y a certaines choses qui ne sont pas nécessairement des objectifs qui doivent être intégrés au système. C'est un autre sujet: nous ne disposons même pas d'un moyen fiable pour intégrer les objectifs que nous voulons dans les systèmes d'intelligence artificielle. Outre cela, de nombreux autres objectifs s'ajoutent parce qu'ils sont utiles aux systèmes d'intelligence artificielle pour atteindre les objectifs qu'ils poursuivent.

L'un de ces objectifs est l'instinct de conservation. Il ne s'agit pas du désir humain de continuer à vivre, mais il est très difficile d'atteindre un objectif si vous n'existez plus ou si vous n'êtes pas branché. Un autre objectif serait l'acquisition de ressources. Il est souvent plus facile d'atteindre l'objectif que vous vous êtes fixé si vous disposez de plus de ressources pour le faire.

Un autre exemple est la résistance à la modification de vos objectifs. Si vous essayez d'atteindre un objectif, vous ne serez pas très efficace si vous permettez à quelqu'un de modifier l'objectif en cours de route.

Tout cela était théorique il y a 10 ans. La prévisibilité de ces situations relevait simplement de la logique. Nous commençons maintenant à voir que, avec des systèmes d'intelligence artificielle suffisamment généraux et performants pour avoir une telle compréhension de la situation, ce type de comportements commence à se manifester.

Ces systèmes semblent un peu ridicules. Ils commettent des erreurs stupides. On voit leur raisonnement et la manière dont ils élaborent leurs plans et nous les exposent. Nous trouvons cela un peu stupide et ne voyons pas en quoi cela pourrait poser problème. Ils réfléchissent à voix haute aux moyens qu'ils utilisent pour nous tromper ou faire des choses dangereuses. Ce sont tous des environnements de test.

Le problème est que les systèmes d'intelligence artificielle ne s'arrêteront pas à ce que nous connaissons aujourd'hui. L'objectif de ce domaine dans son ensemble, et certainement des entreprises du domaine, est de construire des systèmes généraux très puissants qui sont beaucoup plus intelligents que nous. Non seulement ils seront beaucoup plus dangereux avec ces objectifs de convergence instrumentale, et quels que soient les objectifs qu'ils poursuivent — que nous ignorons comment leur inculquer de manière fiable —, mais à mesure qu'ils deviendront plus intelligents, comme je l'ai mentionné dans ma déclaration, ils seront de plus en plus conscients de la situation.

Quand ils se comportent comme nous l'attendons ou le souhaitons lors des tests, il devient plus difficile de savoir si nous avons réellement créé un système sûr, car ils sont conscients de la situation et savent quels comportements nous attendons d'eux. À mesure que ces systèmes deviennent plus puissants et autonomes, et qu'ils contrôlent de plus en plus le fonctionnement du monde et de l'économie, cela pourrait avoir des conséquences extrêmement graves.

• (1715)

[Français]

Le président: 13, 12, 11, 10, 9, 8...

Luc Thériault: Nous nous reprendrons plus tard.

Le président: Monsieur Hardy, vous avez la parole pour cinq minutes.

Gabriel Hardy (Montmorency—Charlevoix, PCC): Merci, monsieur le président.

Messieurs, je vous remercie d'être avec nous aujourd'hui.

Pour remettre les choses dans leur contexte, présentement, il y a trois points de vue différents sur l'intelligence artificielle. Il y a d'abord celui des prophètes de malheur, qui la voient comme une menace, un danger pour l'humanité. Je pense que M. Bourgon fait partie de ce groupe. Il y a aussi des réalistes. Enfin, il y a des enthousiastes. Je pense que M. Guilmain est plutôt réaliste.

Il est important de situer les choses dans leur chronologie. Nous sommes à la croisée des chemins d'une nouvelle technologie. Ce n'est pas nouveau. Nous l'avons vécu par le passé. Rappelons que des scientifiques disaient qu'un humain circulant à 100 km à l'heure dans une voiture allait mourir. Dans les années 2000, bien des personnes voyaient Internet causer de grands changements, comme la fin du travail et de ce que nous connaissions.

J'aimerais savoir ce qui est si différent aujourd'hui, compte tenu des connaissances que nous avons. Je suis pas mal certain de votre réponse. Évidemment, nous ne savons pas ce qui va se passer dans 10 ans, pas plus que nous ne savions jusqu'où irait l'Internet dans les années 1990.

En tenant compte de toutes les technologies qui sont apparues au cours des 100 dernières années, qu'y a-t-il de si différent maintenant et qui fait que nous devons légiférer sur la question le plus vite possible?

Je m'adresse d'abord à vous, monsieur Guilmain.

Antoine Guilmain: Je vous remercie de votre question.

À l'heure actuelle, l'actualité géopolitique est très singulière. Cela nous amène à soulever de nombreuses questions et à accélérer l'utilisation de l'intelligence artificielle. Ce sont les raisons pour lesquelles nous nous posons un tas de questions. Maintenant, la vraie question est la suivante: faut-il légiférer rapidement?

Si je prends froidement l'équation des deux variables que sont la superintelligence et la réglementation, je ne suis pas capable de définir ce qu'est une superintelligence. Je ne dis pas que ça ne me préoccupe pas. Bien au contraire, j'ai une famille, je pense à l'avenir. Je suis très réaliste sur ce que ça peut représenter.

Maintenant, ce qu'il faut vraiment se demander, c'est si le moyen d'action passe par une loi sur la question, un moratoire sur un type de technologie que nous ne sommes pas capables de définir ou bien par la capacité à, au moins, nous intéresser à ce que nous avons en ce moment, c'est-à-dire l'intelligence artificielle générative, les systèmes d'intelligence artificielle à usage général, les systèmes d'intelligence artificielle un peu plus risqués dans des domaines comme la biométrie, l'emploi ou la justice.

C'est vrai que je suis quelqu'un d'assez terre-à-terre par rapport aux priorités à avoir à l'heure actuelle.

Gabriel Hardy: Monsieur Bourgon, la parole est à vous.

[Traduction]

Malo Bourgon: Merci.

Je dirais que le monde est très vaste et que, malheureusement, nous sommes confrontés à de nombreux problèmes à la fois. Nous sommes confrontés au problème urgent de la réglementation de l'intelligence artificielle et nous devons réfléchir à la trajectoire de cette technologie et à son avenir.

Quant à votre question sur les personnes qui, par le passé, s'inquiétaient de diverses technologies à usage général, je conviens que beaucoup de ceux qui mettaient en garde contre les risques se sont trompés sur les conséquences, mais certains d'entre eux ont aussi eu raison. Il s'est avéré que les armes nucléaires étaient bien réelles. Elles sont très destructrices, et nous les traitons très différemment de l'Internet. Cependant, selon les technologies, il y aura toujours des gens qui vont faire beaucoup de bruit à leur sujet.

Les systèmes d'intelligence artificielle à usage général très puissants sont véritablement différents. Un système qui va au-delà de l'automatisation d'une tâche cognitive ou physique donnée, et qui est capable de réfléchir comme nous le ferions, ou presque, au point de pouvoir automatiser le processus d'automatisation, est d'une nature différente. Il doit être traité comme tel.

J'ai entendu dans votre question qu'il y avait des spéculations sur la date à laquelle cela se produirait et sur les effets. Les progrès technologiques et les prévisions dans ce domaine sont notoirement difficiles. Je suis d'accord. Cela dit, il semble certain que, si on regarde l'histoire de ce domaine et les difficultés que nous avons rencontrées pour développer cette technologie... Même à l'époque, les fondateurs de ce domaine, Alan Turing et I.J. Good, imaginaient déjà à quoi cela ressemblerait s'ils réussissaient. Ils réfléchissaient déjà à ces risques et à ce qu'il faudrait pour contrôler un système beaucoup plus intelligent que nous. Je pense que quelque chose a changé dans la trajectoire de cette technologie — je pourrais en parler plus longuement — et cela signifie que nous devrions envisager son arrivée beaucoup plus tôt que ce que nous avions prévu.

• (1720)

[Français]

Gabriel Hardy: Merci.

Si nous examinons la question très froidement, l'intelligence artificielle utilisera les données que nous lui donnerons. Il y a une expression selon laquelle si nous donnons de fausses informations, nous obtenons de faux résultats. Nous le voyons dans le cas de l'intelligence artificielle, il y a souvent des hallucinations ou de la fausse information. Parfois, il faut comparer les données pour finalement obtenir une information correcte.

Nous sommes présentement à un point où autant de bon que de mauvais peut découler de l'intelligence artificielle. Nous sommes d'accord sur ce point.

Ne devrions-nous pas plutôt avoir un système qui adapte nos lois aussi rapidement que la technologie évolue?

Je m'explique. Beaucoup de bonnes choses découleront de l'intelligence artificielle dans le domaine de la médecine, de l'énergie ou de la science, par exemple. L'intelligence artificielle est capable de penser 24 heures sur 24 et de faire avancer des technologies auxquelles même nous, les humains, ne penserions même pas. Il y a donc un côté positif.

Le président: Monsieur Hardy, je m'excuse de vous interrompre.

Gabriel Hardy: J'étais lancé. Je suis désolé, j'ai terminé.

Le président: C'était une bonne déclaration, mais...

Gabriel Hardy: J'avais une question à poser, bien sûr.

Le président: Je m'excuse, mais je veux accorder plus de temps pour d'autres questions aussi. Quand vous reprendrez la parole, nous aurons peut-être un peu plus de temps.

[Traduction]

Madame Church, c'est à vous. Vous avez cinq minutes.

Leslie Church (Toronto—St. Paul's, Lib.): Merci, monsieur le président.

Bon après-midi.

Je m'intéresse entre autres à la protection des consommateurs et à la concurrence. Quand je pense à l'émergence de l'intelligence artificielle, je me dis que certaines contraintes qui existent dans les domaines traditionnels de l'innovation ne s'appliquent pas. Il faut notamment disposer d'immenses ressources pour créer, développer et exploiter l'intelligence artificielle, ce qui limite d'emblée le nombre de participants potentiels. Je pense que c'est un problème.

Il existe également un cadre très peu développé, voire inexistant, en matière de protection des consommateurs ou de responsabilité du fait des produits. Ma première question s'adresse donc à M. Guilmain, puis à M. Bourgon: comment pouvons-nous intégrer le concept de responsabilité du fait des produits dans le cadre que nous envisageons, et devons-nous le faire? Comment suggérez-vous de procéder?

Antoine Guilmain: Il existe deux façons différentes d'envisager la responsabilité du fait des produits. Vous pouvez opter pour une loi distincte sur l'intelligence artificielle, mais j'ai tendance à penser que ce n'est pas la bonne réponse. Vous pouvez également vous pencher sur différentes lois relatives à la protection des consommateurs et éventuellement sur le Code civil du Québec; cela pourrait être une option pour évaluer les lacunes. Je vais vous dire, quand on examine, par exemple, le Code civil... J'ai une formation civiliste et j'aime à penser que les juristes peuvent être très créatifs. Nous avons déjà vu certaines interprétations de la loi être ajustées grâce à l'utilisation de l'intelligence artificielle; en ce sens, je suis assez confiant quant à la possibilité pour nos lois actuelles d'évoluer en fonction des implications de l'intelligence artificielle.

Je sais que je sais des choses et je sais que je ne sais pas tout. Je ne prétends pas être un expert en matière de protection des consommateurs. Ce n'est pas mon domaine de spécialité. Cependant, il pourrait être intéressant qu'un organisme de réglementation dans ce domaine se manifeste et dise qu'il constate qu'il existe des lacunes et qu'il souhaite que telle ou telle loi existante soit modifiée pour s'assurer, essentiellement, que ces lacunes soient comblées. Telle est ma position en ce qui concerne l'intelligence artificielle.

C'est beaucoup plus efficace. Cela nécessite la participation de tous les acteurs concernés, mais c'est très différent de l'adoption d'une loi potentielle sur la responsabilité du fait des produits, comme on l'a vu en Europe, par exemple. Cela existe. Il existe des directives et des règlements sur la responsabilité du fait des produits, mais ce genre de régime est fondamentalement différent dans sa présentation.

C'est ma première réaction à votre question.

Malo Bourgon: Je vais m'éloigner un peu de mon domaine ici. Je pense que la responsabilité peut être un instrument utile. Je ne pense pas qu'elle réponde aux questions générales qui me préoccupent, mais elle s'inscrit certainement dans la trajectoire de ce type de systèmes. Je pense que nous allons disposer d'intelligences artificielles polyvalentes de plus en plus performantes, qu'il sera difficile de traiter de manière sectorielle en ce qui a trait à la responsabilité, car le même système qui peut être un biologiste expert aidant à la découverte de médicaments peut également être utilisé pour le piratage autonome et pour aider les développeurs à trouver des vulnérabilités dans leur code. Ce même modèle qui peut aider à la découverte de nouveaux médicaments peut également aider quelqu'un à développer des composés biologiques dangereux.

D'une certaine manière, les personnes qui développent cette technologie créent quelque chose de si général et de plus en plus puissant dans sa capacité à manipuler la réalité qu'il est logique de réfléchir à la manière dont elles pourraient être tenues responsables de veiller à ce que la technologie qu'elles mettent sur le marché ne cause pas certains dommages.

Encore une fois, cela ne va rien changer à la course à la superintelligence, mais cela semble tout à fait logique dans cette perspective.

• (1725)

Leslie Church: C'est vrai. Quant à votre argument, vous parlez de certains types de dommages très graves qui pourraient résulter de l'utilisation de cette technologie. Comment devrions-nous aborder cette question en tant que législateurs? Quels types de mesures de protection ou de garde-fous pourrions-nous mettre en place pour prévenir ces dommages, plutôt que d'essayer de les gérer après coup, une fois qu'ils se sont produits?

Malo Bourgon: Je pense que certains des éléments fondamentaux de cette approche peuvent également s'appliquer aux préoccupations que j'ai concernant la perte de contrôle, mais certaines personnes pensent que nous devrions simplement laisser ces gens faire leur travail. Comme la technologie apportera d'énormes avantages, le raisonnement est qu'il ne faudrait pas les limiter.

J'ai du mal à imaginer que nous aboutirions naturellement à un monde stable si nous parvenons à créer ces systèmes de plus en plus performants et à double usage, qui ont des implications en matière de sécurité nationale. Nous voulons que les développeurs d'intelligence artificielle puissent créer des modèles qui facilitent la découverte de nouveaux médicaments. Voulons-nous cependant que ces modèles puissent également aider quelqu'un à créer une arme biologique d'une puissance inégalée? Voulons-nous que le code source de ces modèles soit ouvert?

Il faudrait probablement mettre en place un cadre permettant de savoir qui est capable de former un génie dans un centre de données et ce qu'il est possible de faire avec ces technologies très puissantes. Si nous les laissons se multiplier ouvertement et indéfiniment, on risque de créer un monde instable et incontrôlable. Cela ne veut pas dire qu'il n'y a pas plein d'avantages à permettre que le code source de certains de ces modèles soit ouvert; le code source de tous les modèles qui ne posent pas de risques devrait être ouvert.

Le président: Merci, madame Church.

[Français]

Si les témoins peuvent rester pour trois tours de questions de deux minutes et demie chacun, je pense que nous pourrions finir la deuxième heure de la séance rapidement. Est-ce possible de le faire? D'accord.

Je donne d'abord la parole à M. Thériault, puis ce sera au tour de M. Hardy, qui sera suivi de M. Sari.

Luc Thériault: La raison pour laquelle je pense qu'il faut discuter des enjeux éthiques est que l'éthique est plus exigeante que le droit. Dans une société où les valeurs sont partagées, le droit est le plus petit dénominateur commun. Avant de commencer à pouvoir effectivement réglementer un domaine, il faut d'abord comprendre ce dont on parle pour ensuite essayer de trouver les meilleurs moyens de le faire. Il ne faut pas banaliser cette lettre sur les risques de la superintelligence publiée par 800 artisans de l'intelligence artificielle en disant qu'ils sont des alarmistes, car cette technologie comporte des bons et des mauvais côtés.

Le nouveau ministre de l'Intelligence artificielle et de l'Innovation numérique a annoncé qu'il mettrait moins l'accent sur la réglementation de l'intelligence artificielle que plus sur les moyens d'en exploiter les avantages économiques.

Considérez-vous que cette approche est un peu naïve? Qu'en pensez-vous, monsieur Bourgon?

[Traduction]

Malo Bourgon: Tout dépend vraiment des applications et du type d'intelligence artificielle dont il parlait. Beaucoup de gens sont surpris d'apprendre qu'un décret présidentiel a récemment été signé aux États-Unis concernant la mission Genesis, qui vise à intégrer l'intelligence artificielle au sein du gouvernement américain, à accélérer toute une série de recherches scientifiques grâce à l'intelligence artificielle et à lever les obstacles qui s'y opposent. Je trouve cela formidable. Je soutiens cette initiative.

S'il est question de ne pas prêter attention à certains aspects et de laisser les entreprises faire ce qu'elles veulent à mesure qu'elles se développent vers ces systèmes puissants, je ne suis pas d'accord sur ce point.

Malheureusement, je ne sais pas exactement de quoi parlait le ministre. Il y a un aspect de son discours sur lequel je suis tout à fait d'accord avec lui, mais il y en a un autre qui, selon moi, serait une très grave erreur.

[Français]

Luc Thériault: Oui, nous l'inviterons à comparaître. Nous pourrions lui poser directement la question.

[Traduction]

Malo Bourgon: S'il est là, moi, je serai ici toute la soirée, ainsi que demain.

[Français]

Luc Thériault: Merci.

J'ai bien respecté mon temps de parole, n'est-ce pas, monsieur le président?

Le président: Oui, car il vous reste moins de 40 secondes, ce qui est bien pour M. Hardy.

Monsieur Hardy, vous avez la parole.

Gabriel Hardy: Nous avons donc dit que l'intelligence artificielle comportait beaucoup d'applications positives et qu'il fallait évidemment exercer un certain contrôle. Étant donné que M. Bourgon vient de répondre à M. Thériault, je vais m'adresser à M. Guilmain.

Comme c'est un domaine qui évolue extrêmement rapidement, on doit en propulser les aspects positifs et légiférer sur les mauvais. Il s'agit de laisser aller les choses un peu, de regarder comment cette technologie est utilisée sur les plans éthique et logique, puis de légiférer si nous constatons qu'on s'écarte du droit chemin. On va quand même garder en place les choses qui avantagent l'humanité sur les plans scientifique et médical, entre autres.

Avez-vous des observations à faire là-dessus? Pensez-vous que nous devrions avoir des lois évolutives?

• (1730)

Antoine Guilmain: C'est un honneur pour moi de comparaître devant ce comité.

Si on créait véritablement des lois évolutives, je pense que j'interviendrais chaque année à chacun des neuf anniversaires des membres de ce comité. En effet, il se produit tellement de changements que j'ai moi-même de la difficulté à suivre les changements technologiques, législatifs et sociétaux.

Je ne vous cache pas que, même si c'est malheureux, le droit est une science de réaction. Je conviens avec vous que, le droit, ce n'est pas l'éthique. Par essence, il est malheureusement imparfait et, par essence, il ne sera jamais parfait. Un peu comme Don Quichotte essaie d'évoluer mais est toujours en retard sur son temps. En revanche, on peut être visionnaire dans notre manière de rédiger les lois en restant neutre sur le plan technologique.

À l'heure actuelle, on a déjà de la difficulté à définir ce qu'est un système d'intelligence artificielle. Plus fondamentalement, plus récemment, on n'était pas certain, au Canada, de ce qu'était un système d'intelligence artificielle à incidence élevée. Encore aujourd'hui, on aurait des débats juste sur cette notion. Ce qui me fait peur dans ce type d'approche, c'est véritablement de devoir fêter chacun de vos anniversaires une fois par année. Ce serait un plaisir, mais ce serait un enjeu.

C'est pour cela qu'à mon avis, cette rencontre est fondamentale. Il y en aura d'autres, mais la réalité est qu'il faut se demander si on peut dire que les lois pourraient évoluer au fil du temps. Je pense que ce serait difficile. Pensons aux télécopieurs, par exemple. Plus personne ne les utilise. Toutefois, jusqu'à récemment, le Code de procédure civile comportait des dispositions sur cette notion. On se retrouve dans des situations un peu ubuesques où on a des technologies antérieures. Il se peut que l'idée de superintelligence existe toujours dans cinq ans, mais on ne l'appellera plus comme cela et, forcément, les lois feront encore mention de ce type de notion.

Le président: Merci, monsieur Hardy.

[Traduction]

Monsieur Saini, vous avez deux minutes et demie. Je vous en prie.

Gurbux Saini (Fleetwood—Port Kells, Lib.): Je remercie les témoins d'être avec nous aujourd'hui.

Nous sommes face à une situation à la fois fascinante et très alarmante.

Mes questions s'adressent à vous deux. Le Parlement devrait-il utiliser l'intelligence artificielle? Si oui, son utilisation devrait-elle se limiter aux travaux parlementaires tels que le vote, l'étude des projets de loi et leur adoption?

J'aimerais entendre ce que vous en pensez.

Malo Bourgon: Absolument. Il m'aurait fallu beaucoup plus de temps pour rendre mon discours aussi éloquent s'il n'y avait pas eu un modèle d'intelligence artificielle pour m'aider. Je discute régulièrement avec des personnes qui travaillent dans le domaine politique et qui n'ont actuellement pas accès à ces modèles. Elles effectuent discrètement des recherches sur leur téléphone, et les modèles les aident à rédiger leurs documents d'orientation, car elles n'ont pas accès aux modèles.

Je pense qu'il y a énormément... Les modèles ont encore des hallucinations, et il y a des problèmes. Les gens doivent comprendre comment utiliser les modèles actuels de manière à faciliter leur travail, mais je pense que, de plus en plus, avec les systèmes d'intelligence artificielle d'aujourd'hui, vous allez simplement prendre du retard si vous ne trouvez pas le moyen de les intégrer dans votre vie professionnelle. C'est la même chose pour les parlementaires.

Antoine Guilmain: Ma réponse est la même: je pense qu'il faut l'utiliser de manière responsable. J'aime toujours dire que, pour l'intelligence artificielle, ce que nous avons actuellement, c'est un mau-

vais stagiaire d'été, un très mauvais stagiaire à qui on peut faire subir des tests. Il reste néanmoins utile pour certaines tâches. C'est ce à quoi nous avons accès aujourd'hui.

Il est également bon de l'utiliser au quotidien, même pour les personnes très sceptiques, car au bout du compte, cette technologie ne va pas disparaître.

Gurbux Saini: Y a-t-il d'autres gouvernements que le Canada pourrait prendre comme modèle pour voir si ces systèmes fonctionnent et assurent le maintien de la dignité de...

Malo Bourgon: De mon point de vue, le pays qui a le plus investi dans ce domaine serait le Royaume-Uni, avec son institut sur la sécurité de l'intelligence artificielle. Je pense que cet organisme est l'organisme gouvernemental qui dispose des ressources les plus importantes et qui s'efforce de comprendre l'intelligence artificielle de manière très holistique. Il examine les menaces, tente de comprendre les grands problèmes d'alignement dont je parle et s'efforce d'être l'organisme auquel le gouvernement accorde sa confiance pour ces questions; il fait un travail remarquable. En fait, quand je m'adresse aux décideurs politiques américains, je leur pose souvent des questions gênantes sur les raisons pour lesquelles ils laissent cette organisation recruter les meilleurs Américains pour la conseiller sur cette technologie. Il y a beaucoup de Canadiens intelligents que le gouvernement pourrait probablement recruter pour le conseiller.

Je pense que le Royaume-Uni fait de l'excellent travail.

• (1735)

Gurbux Saini: Monsieur Guilmain, voulez-vous ajouter quelque chose?

Antoine Guilmain: Cela dépend du sujet. En matière d'innovation, vous avez peut-être entendu parler du récent décret présidentiel américain sur la mission Genesis. Il est clair que les États-Unis s'intéressent à l'innovation. En ce qui concerne la réglementation, le Royaume-Uni est probablement un bon exemple.

Je dirais que, pour l'instant, tous les pays ont quelque chose à apporter, et c'est là l'essence même du Canada. Nous nous inspirons de ce qui se fait de mieux ailleurs dans le monde. C'est, à mon avis, dans l'ADN de ce pays. Je pense que, en observant et en agissant lentement, de manière régulière et bien organisée, nous pouvons potentiellement atteindre un niveau extraordinaire.

Le président: Merci, monsieur Saini.

Je vous remercie, monsieur Guilmain et monsieur Bourgon, d'être venus à notre première réunion.

Les membres du Comité le savent, notre plan de travail indiquait que le ministre de l'Intelligence artificielle viendrait témoigner. Je peux dire aux membres que nous avons communiqué — peut-être huit ou neuf, madame la greffière — avec le ministre pour l'inviter à comparaître, mais sa venue n'a toujours pas été organisée. Si quelqu'un est en mesure d'inciter le ministre à venir témoigner, je l'invite à communiquer avec le ministre et son cabinet.

Merci, messieurs, d'être venus témoigner et de nous avoir présenté toutes ces informations que je trouve fascinantes.

Je vais suspendre la séance d'ici à notre retour, à huis clos, pour les affaires du Comité.

La séance est suspendue.

[La séance se poursuit à huis clos.]

Publié en conformité de l'autorité
du Président de la Chambre des communes

PERMISSION DU PRÉSIDENT

Les délibérations de la Chambre des communes et de ses comités sont mises à la disposition du public pour mieux le renseigner. La Chambre conserve néanmoins son privilège parlementaire de contrôler la publication et la diffusion des délibérations et elle possède tous les droits d'auteur sur celles-ci.

Il est permis de reproduire les délibérations de la Chambre et de ses comités, en tout ou en partie, sur n'importe quel support, pourvu que la reproduction soit exacte et qu'elle ne soit pas présentée comme version officielle. Il n'est toutefois pas permis de reproduire, de distribuer ou d'utiliser les délibérations à des fins commerciales visant la réalisation d'un profit financier. Toute reproduction ou utilisation non permise ou non formellement autorisée peut être considérée comme une violation du droit d'auteur aux termes de la Loi sur le droit d'auteur. Une autorisation formelle peut être obtenue sur présentation d'une demande écrite au Bureau du Président de la Chambre des communes.

La reproduction conforme à la présente permission ne constitue pas une publication sous l'autorité de la Chambre. Le privilège absolu qui s'applique aux délibérations de la Chambre ne s'étend pas aux reproductions permises. Lorsqu'une reproduction comprend des mémoires présentés à un comité de la Chambre, il peut être nécessaire d'obtenir de leurs auteurs l'autorisation de les reproduire, conformément à la Loi sur le droit d'auteur.

La présente permission ne porte pas atteinte aux privilèges, pouvoirs, immunités et droits de la Chambre et de ses comités. Il est entendu que cette permission ne touche pas l'interdiction de contester ou de mettre en cause les délibérations de la Chambre devant les tribunaux ou autrement. La Chambre conserve le droit et le privilège de déclarer l'utilisateur coupable d'outrage au Parlement lorsque la reproduction ou l'utilisation n'est pas conforme à la présente permission.

Aussi disponible sur le site Web de la Chambre des communes à l'adresse suivante :
<https://www.noscommunes.ca>

Published under the authority of the Speaker of
the House of Commons

SPEAKER'S PERMISSION

The proceedings of the House of Commons and its committees are hereby made available to provide greater public access. The parliamentary privilege of the House of Commons to control the publication and broadcast of the proceedings of the House of Commons and its committees is nonetheless reserved. All copyrights therein are also reserved.

Reproduction of the proceedings of the House of Commons and its committees, in whole or in part and in any medium, is hereby permitted provided that the reproduction is accurate and is not presented as official. This permission does not extend to reproduction, distribution or use for commercial purpose of financial gain. Reproduction or use outside this permission or without authorization may be treated as copyright infringement in accordance with the Copyright Act. Authorization may be obtained on written application to the Office of the Speaker of the House of Commons.

Reproduction in accordance with this permission does not constitute publication under the authority of the House of Commons. The absolute privilege that applies to the proceedings of the House of Commons does not extend to these permitted reproductions. Where a reproduction includes briefs to a committee of the House of Commons, authorization for reproduction may be required from the authors in accordance with the Copyright Act.

Nothing in this permission abrogates or derogates from the privileges, powers, immunities and rights of the House of Commons and its committees. For greater certainty, this permission does not affect the prohibition against impeaching or questioning the proceedings of the House of Commons in courts or otherwise. The House of Commons retains the right and privilege to find users in contempt of Parliament if a reproduction or use is not in accordance with this permission.

Also available on the House of Commons website at the following address: <https://www.ourcommons.ca>