



CHAMBRE DES COMMUNES
HOUSE OF COMMONS
CANADA

45^e LÉGISLATURE, 1^{re} SESSION

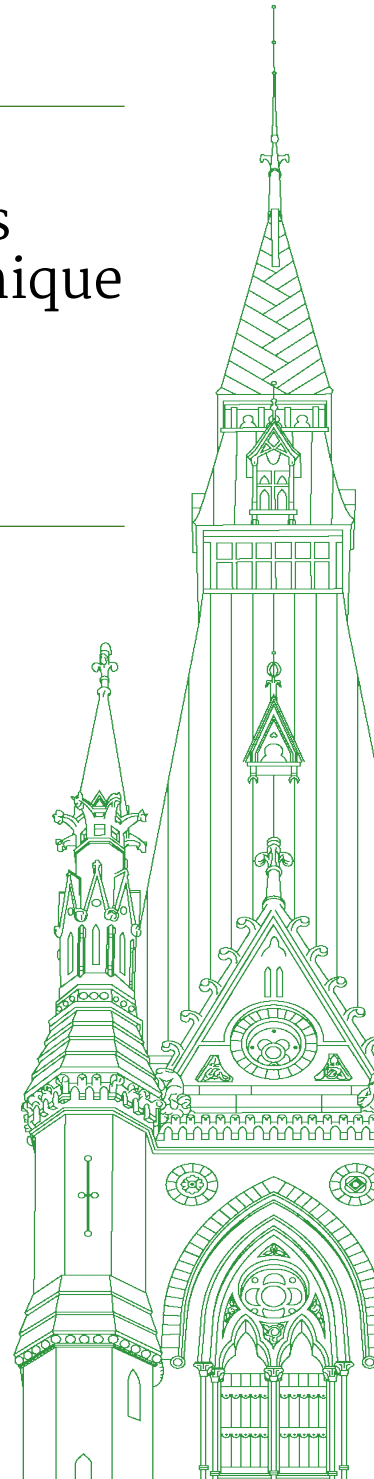
Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique

TÉMOIGNAGES

NUMÉRO 020

Le lundi 1^{er} décembre 2025

Président : John Brassard



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

Le lundi 1er décembre 2025

• (1100)

[Traduction]

Le président (John Brassard (Barrie-Sud—Innisfil, PCC)):
Bonjour à tous. Nous voilà en décembre, et je déclare la séance ouverte.

Je vous souhaite à toutes et à tous la bienvenue à la 20^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 reprend son étude sur les défis que posent l'intelligence artificielle et son encadrement.

[Traduction]

Je suis heureux d'accueillir nos deux témoins pour cette première heure, qui représentent Conjecture Ltd. Il s'agit de Connor Leahy, directeur général, et de Gabriel Alfour, directeur technique.

Monsieur Leahy, vous avez un maximum de cinq minutes pour vous adresser au Comité. Je crois comprendre que vous aurez peut-être besoin d'un peu plus de temps. Si cela vous prend jusqu'à six minutes, je vous laisserai poursuivre, mais je sais que nous avons beaucoup de questions à vous poser.

Monsieur Leahy, vous avez la parole.

Connor Leahy (directeur général, Conjecture Ltd.): Merci, monsieur le président et distingués membres du Comité, de m'avoir invité à témoigner aujourd'hui.

Je suis expert des menaces d'ampleur catastrophique que l'IA fait peser sur le monde et c'est principalement dans cette perspective que je m'adresserai à vous.

Je suis le PDG de Conjecture, une entreprise de recherche sur la sécurité de l'IA, l'intelligence artificielle. Je suis également conseiller à ControlAI, un organisme sans but lucratif qui se concentre sur l'atténuation des risques que l'IA avancée pose pour la sécurité.

En 1985, l'humanité s'est rendu compte qu'il y avait un trou dans le ciel. Les scientifiques ont découvert que les chlorofluorocarbures, les CFC, appauvrissaient la couche d'ozone qui protège normalement l'humanité des dommages causés par le rayonnement ultraviolet. À cette époque, l'humanité vivait une fracture profonde, celle de la guerre froide entre les États-Unis et l'URSS qui la menaçait d'anéantissement nucléaire.

Dans un contexte de fortes tensions géopolitiques, les deux superpuissances ont fini par se serrer la main en signant à la fois un

traité historique de désescalade nucléaire et le Protocole de Montréal, en 1987, pour interdire et éliminer progressivement les CFC. Ce protocole a finalement été ratifié à l'unanimité. Malgré les divisions dans le monde, ces puissances rivales ont fait front commun pour refermer le trou dans l'atmosphère et reconnaître que l'escalade nucléaire sans fin n'était dans l'intérêt de personne; le reste du monde a emboîté le pas.

En 2023, l'humanité a reçu un nouvel avertissement des chercheurs en IA lauréats d'un prix Nobel et des patrons de grandes entreprises d'IA: « L'atténuation du risque d'extinction associé à l'IA doit être une priorité mondiale au même titre que d'autres risques sociétaux tels que les pandémies et la guerre nucléaire. » Ce risque d'extinction est posé par la superintelligence, ce sous-produit de l'IA que les grandes entreprises d'IA cherchent précisément à développer.

On définit la superintelligence comme étant une intelligence artificielle plus compétente que n'importe quel humain, dans toutes les tâches cognitives pertinentes et dans tous les domaines pertinents, capable en outre de fonctionner en l'absence de surveillance et de contrôle humains. S'il devait exister des systèmes autonomes aptes à supplanter n'importe quel être humain dans toutes les tâches pertinentes de la science, des affaires, de la persuasion, de la politique et de la guerre, et si nous ne contrôlions pas ces systèmes, il deviendrait difficile d'imaginer un avenir favorable pour l'humanité.

Le risque découle en grande partie du fait que les développeurs d'IA ne comprennent pas comment fonctionnent réellement les systèmes d'IA qu'ils créent, mais qu'ils ne parviennent pas à développer de façon sécuritaire. Dario Amodei, chef de la direction de la deuxième entreprise d'IA en importance dans le monde, a récemment déclaré que nous « comprenons peut-être 3 % du fonctionnement de l'IA », ce qui, à mon avis, est une surestimation.

Les intelligences artificielles ne sont pas élaborées à coups de codes écrits ligne par ligne, comme c'est le cas pour les logiciels traditionnels. Au lieu de cela, les chercheurs développent essentiellement des modèles d'IA qu'ils nourrissent d'énormes quantités de données et qu'ils forment grâce à des ordinateurs surpuissants pour produire ce qu'on appelle un réseau neuronal plutôt qu'un ensemble de lignes de code informatique.

Malheureusement, le paradigme actuel de développement de l'IA ne permet pas d'appliquer, dès le stade de la conception, les mêmes approches de sécurité que pour d'autres technologies avancées et très risquées. Par exemple, nous ne construirions pas de centrales nucléaires si nous ne savions pas comment contrôler les réactions nucléaires. Les méthodes de contrôle technique accusent un retard considérable par rapport aux progrès réalisés dans les capacités des systèmes d'IA. Il n'existe actuellement aucun règlement juridique contraignant en matière de sécurité de l'intelligence artificielle pour protéger les consommateurs et l'humanité dans son ensemble.

Où en sommes-nous aujourd'hui? À l'heure actuelle, de nombreuses entreprises d'IA investissent des centaines de milliards de dollars pour développer le plus rapidement possible une IA superintelligente, malgré les risques mis en exergue par les experts. Je suis d'avis que cette précipitation vise à déjouer d'éventuelles législations et à leur permettre de mener à bien leurs projets avant que le grand public et les gouvernements ne prennent pleinement conscience des risques tout à fait inadmissibles auxquels le public non consentant est exposé par des acteurs privés imprudents échappant à toute surveillance.

Récemment, les entreprises d'IA se sont lancées dans une course pour automatiser la recherche en IA afin de permettre aux IA de créer elles-mêmes d'autres IA encore plus performantes et d'atteindre ainsi plus rapidement le niveau de superintelligence. On parle alors d'auto-amélioration récursive, ce qui revient à dire que, dès qu'une IA est assez bonne pour en produire de meilleures, c'est qu'il est peut-être déjà trop tard.

Des scientifiques de premier plan estiment maintenant que la superintelligence pourrait être développée d'ici 2030, voire plus tôt. Face à cette menace pressante de la superintelligence, j'aimerais adresser trois recommandations au Comité sur la façon dont le Canada pourrait réagir tout de suite.

Premièrement, le gouvernement canadien devrait reconnaître publiquement la superintelligence comme une menace à la sécurité nationale et mondiale à cause du risque d'extinction qu'elle fait courir à l'humanité.

Deuxièmement, le Canada devrait entamer des négociations en vue de parvenir à un accord international qui consisterait à interdire le développement de la superintelligence, étant donné qu'il n'est pas envisageable de dégager un consensus scientifique qui ne menacerait pas l'humanité d'extinction. L'accord devrait également restreindre et surveiller les précurseurs de la superintelligence, comme l'auto-amélioration récursive.

• (1105)

Troisièmement, le Canada devrait empêcher le développement d'une superintelligence artificielle sur son sol, car celle-ci pourrait dominer les personnes, les entreprises et même l'appareil de sécurité nationale.

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

Le président: Merci, monsieur Leahy.

Monsieur Alfour, vous avez un maximum de cinq minutes pour des remarques liminaires. Allez-y, je vous prie.

Gabriel Alfour (directeur technique, Conjecture Ltd.): Monsieur le président et distingués membres du Comité, je m'appelle Gabriel Alfour. Je suis directeur technique et cofondateur de

Conjecture, une entreprise de recherche sur l'intelligence artificielle. J'ai également aidé à fonder ControlAI, un organisme sans but lucratif qui se consacre à la prévention des risques que l'IA présente pour l'humanité. ControlAI s'est adressé à des législateurs du Canada, des États-Unis, du Royaume-Uni et de l'Union européenne.

Avec l'IA, nous sommes confrontés à de nombreux défis complexes et de taille, mais je pense personnellement et professionnellement que le risque le plus immédiat est celui d'extinction que pose toute IA superintelligente. Les systèmes dotés d'une telle IA dépasseraient largement les capacités cognitives humaines et seraient capables de nous concurrencer sur les plans du développement scientifique et militaire, de la persuasion, de la politique, des affaires et plus encore. Ces produits seraient plus efficaces non seulement pour les particuliers, mais aussi pour les grandes entreprises, les organismes de sécurité nationale et les gouvernements. De telles IA, si elles sont mises au point, décideront de l'avenir à notre place.

Comment en sommes-nous arrivés là avec l'IA?

Tout d'abord, les plus grands experts du domaine — les chercheurs en IA les plus cités et les dirigeants des principaux laboratoires d'IA — ont prévenu en 2023 que « l'atténuation du risque d'extinction lié à l'IA devrait être une priorité mondiale au même titre que les autres risques sociétaux, comme les pandémies et la guerre nucléaire ». Cependant, ces avertissements ont été ignorés. Des entreprises de premier plan dans le domaine de l'IA continuent de chercher, sans précaution aucune, à mettre en place des systèmes d'intelligence artificielle sophistiqués capables de surpasser nos meilleurs experts en technologie, en génie et en sécurité nationale, et de survivre à des tentatives de mise hors circuit. Leurs plans pour contrôler les systèmes superintelligents, quand ils existent, sont au mieux non équilibrés et spéculatifs.

Deuxièmement, une fausse idée a cours au sujet du développement de l'IA, celle que la façon dont ces systèmes se comportent est planifiée, mais tel n'est pas le cas. Nous avons planifié le développement de l'IA jusqu'à il y a une quinzaine d'années, mais de nos jours, les systèmes d'IA modernes sont nourris, plutôt que d'être bâtis, les systèmes d'IA modernes sont nourris par des quantités massives de données. Nous ne pouvons ensuite ni prédire ni contrôler le comportement résultant de ces IA. Autrement dit, l'intelligence artificielle n'est pas codée ligne par ligne par des humains, et les chercheurs et les ingénieurs n'ont pas besoin de la comprendre dans le détail pour la faire évoluer. Quand un système d'IA encourage un jeune à se suicider, qu'il manipule ses utilisateurs ou qu'il résiste à sa mise hors circuit, rien de cela n'est le résultat d'une programmation humaine. C'est la conséquence de ne pas savoir comment diagnostiquer ce qui a amené le système à agir de telle ou telle façon, ou de ne pas être en mesure de l'empêcher de reproduire des comportements.

Enfin, à ce jour, l'intelligence artificielle demeure l'exception et non la norme en matière de réglementation des industries à risque élevé. Pour exercer leurs activités dans des domaines comme le nucléaire et la biotechnologie, les développeurs doivent respecter des normes de sûreté rigoureuses, mettre en œuvre des stratégies d'atténuation des risques, se soumettre à des inspections, etc., mais le domaine de l'intelligence artificielle demeure largement non réglementé malgré les préoccupations croissantes de l'industrie. Des ingénieurs en IA de San Francisco m'ont dit qu'ils ne comprennent pas ce qu'ils bâtissent, et certains considèrent même que ce qu'ils font est clairement dangereux. Même Geoffrey Hinton, le « parrain de l'IA », a quitté Google en lançant une mise en garde contre les risques liés à l'IA.

Que peut-on faire pour prévenir les risques d'extinction que présente la superintelligence artificielle? Je crois que les pays ne devraient pas agir unilatéralement contre leurs propres intérêts, encore moins sur la base d'une confiance aveugle. Ils doivent plutôt faire deux choses.

Premièrement, sur un plan national, ils doivent arrêter le développement des systèmes d'intelligence artificielle les plus dangereux, à savoir les systèmes superintelligents. Chaque pays a tout à perdre du développement de la superintelligence et tout à gagner à suspendre au niveau national tous les programmes visant à la développer. Une fois déployés, ces systèmes ne pourraient pas être arrêtés et seraient plus performants que n'importe quel humain en piratage, entre autres, ce qui menacerait la sécurité nationale des pays.

Deuxièmement, les pays doivent s'entendre entre eux pour réglementer et surveiller les précurseurs de la superintelligence. Nous devrions appliquer la même approche réglementaire que celle utilisée pour les technologies à double usage, comme les matières nucléaires, biologiques et chimiques, et interdire les programmes de développement susceptibles d'occasionner des dommages extrêmes — dans ce cas-ci, une superintelligence artificielle — tout en réglementant leurs précurseurs. Cela permettra aux applications bénéfiques de prospérer tout en prévenant les dommages catastrophiques.

La détermination des capacités de réglementation des précurseurs est une cible mouvante qui évoluera en même temps que notre compréhension de l'IA. Certains précurseurs sont malheureusement à double usage. Les ordinateurs et les centres de données sont avantageux sur le plan économique, mais essentiels au développement d'une superintelligence. De même, les capacités de piratage offrent des avantages militaires, mais pourraient permettre à l'IA de passer outre les protections cybernétiques. Dans le cas de tels précurseurs à double usage, les accords internationaux sont essentiels. Aucun pays ne peut, à lui seul, atténuer ce genre de risques, pas plus qu'un pays ne devrait assumer le coût économique d'éventuelles restrictions tandis que d'autres iraient de l'avant.

• (1110)

Entretemps, certains précurseurs ont des applications plus restreintes qui se limitent à la recherche sur l'IA proprement dite, comme des systèmes capables de faire progresser de façon autonome la recherche sur l'IA sans surveillance humaine, ce qui pourrait déclencher une boucle de rétroaction non contrôlée d'améliorations des capacités.

Le Canada peut également agir à l'échelle nationale pour neutraliser les systèmes d'IA dangereux à l'intérieur de ses frontières. Par exemple...

Le président: Monsieur Alfour, je vais devoir vous arrêter ici, car je sais que les membres du Comité veulent passer aux questions. Vous pourrez peut-être nous faire profiter du reste de votre déclaration dans vos réponses.

Je veux aussi m'assurer que vous puissiez tous les deux vous exprimer dans la langue de votre choix, parce qu'il y aura des questions en anglais et en français.

[Français]

Nous allons commencer en français.

Monsieur Hardy, vous avez la parole pour six minutes.

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

Messieurs, merci d'être des nôtres aujourd'hui.

C'est un sujet extrêmement important. Nous n'en sommes qu'aux débuts de l'étude, mais les témoins qui viennent nous parler de l'intelligence artificielle semblent avoir des opinions assez vastes. Il y en a qui nous parlent du côté très risqué et très dangereux de l'intelligence artificielle, et il y en a d'autres qui sont très positifs envers ce qu'elle va nous apporter.

J'aimerais qu'on fasse le parallèle entre l'intelligence artificielle et ce qu'on voit concernant les réseaux sociaux. On a laissé les entreprises privées développer directement les réseaux sociaux à une vitesse folle avec, en quelque sorte, un bébé d'intelligence artificielle qui analyse tout ce que nous regardons afin d'essayer de tout le temps garder notre attention sur ces réseaux. Nos jeunes n'ont jamais été aussi stressés, notamment à cause de ça.

Quelle comparaison feriez-vous entre le début de l'intelligence artificielle sur les réseaux sociaux et la superintelligence qui est en train de se développer, soit celle que vous nous présentez depuis tantôt?

• (1115)

[Traduction]

Connor Leahy: Je vais répondre à cette question.

Une grande partie des premières recherches qui ont conduit à l'essor actuel de l'intelligence artificielle ont débuté dans le contexte des réseaux sociaux. Une grande partie des premières recherches sur ce que l'on appelle aujourd'hui l'apprentissage profond et l'IA ont été menées pour les algorithmes de recommandation des médias sociaux.

Personnellement, je fais partie d'une génération plus jeune, à la queue de la génération Z. Je me souviens qu'on nous avait promis, en quelque sorte, que si on laissait les réseaux sociaux et Internet fonctionner librement, sans réglementation, cela apporterait liberté et prospérité au monde. Certains députés se souviennent peut-être du printemps arabe. À l'époque, beaucoup, comme moi, pensaient que l'accès généralisé à Internet et aux médias sociaux nous apporterait la liberté et la démocratie.

Ces promesses se sont avérées être des mensonges. Elles ne se sont pas concrétisées. Elles sont plutôt utilisées par les entreprises de médias sociaux pour cannibaliser de nombreux aspects de nos communications interpersonnelles et de nos relations interpersonnelles dans leur propre intérêt. Elles dépensent maintenant des centaines de millions de dollars en lobbying sur diverses tribunes pour essayer d'empêcher les gens de s'ingérer.

Je pense que le schéma comportemental — consistant à développer une technologie si rapidement que les gouvernements ne sont pas en mesure de réagir, et tenter activement de freiner l'action gouvernementale pour empêcher toute réglementation de cette technologie jusqu'à ce qu'il soit trop tard — est exactement le même, d'où la stratégie que les entreprises d'IA sont en train de déployer.

[Français]

Gabriel Hardy: Si je comprends bien, au fond, on est un peu en train de répéter l'erreur qu'on a faite par le passé avec les réseaux sociaux. Ce sont les incitatifs qui font que les compagnies ont tout intérêt à développer le plus vite possible une intelligence artificielle très forte, c'est-à-dire la superintelligence artificielle, pour prendre le dessus sur le marché.

Dans votre allocution d'ouverture, vous avez fait la comparaison entre l'évolution du nucléaire et le Protocole de Montréal. Selon vous, quelles seraient les choses à faire pour que les gouvernements interviennent le plus vite possible en matière d'intelligence artificielle? Que pourrait-on faire pour que l'on comprenne les dangers et les côtés positifs de l'intelligence artificielle, pour qu'on intervienne le plus vite possible pour garder l'humain à l'intérieur du mécanisme et pour qu'on s'assure que l'intelligence artificielle se développe seulement pour son côté positif et que son côté négatif peut être contrôlé?

[Traduction]

Connor Leahy: C'est tout à fait exact. Je pense que nous com-mettons la même erreur dans une certaine mesure, et qu'il faudrait éviter de tomber dans ce piège. Par conséquent, il est extrêmement important que le gouvernement agisse rapidement. Comme je l'ai dit dans mes propos liminaires, le plus important est de stopper car-rément le développement de l'IA superintelligente et dangereuse.

Il ne s'agit pas d'un pan ésotérique et marginal de l'industrie. C'est un objectif que les services de recherche de ces entreprises annoncent comme leur priorité dans leurs publicités. Leurs représentants se présentent à des réceptions où ils se vantent de mettre au point une superintelligence. La manœuvre n'a rien de secret. Il s'agit d'une initiative à grande échelle, et le gouvernement pourrait, dès maintenant, agir à cet égard. Déjà, le simple fait de reconnaître ces risques et de les intégrer dans le discours aux échelons national et international constituerait une première étape sur la voie de la stigmatisation et, éventuellement, de la neutralisation de cette dangereuse évolution, tout en mettant la table pour des négociations sur la façon de traiter les précurseurs à double usage.

C'est un défi très difficile à relever sur le plan de la réglementation. C'est pourquoi des échanges comme ceux d'aujourd'hui sont si importants. Personnellement, j'estime que la première étape doit consister à empêcher la création d'une superintelligence à l'échelle nationale et internationale. Cela veut dire qu'il faudra s'orienter vers une réglementation raisonnable des technologies à double usage et tirer parti des leçons qui ont été apprises dans d'autres domaines technologiques à haut risque.

[Français]

Gabriel Hardy: Merci.

J'aimerais aussi entendre vos commentaires sur le sujet, monsieur Alfou. Dans votre présentation, vous avez parlé des risques. Or, je crois qu'il y a un autre risque que vous n'avez pas mentionné. Vous disiez que le Canada devrait cesser les études visant à développer une superintelligence artificielle. Toutefois, si le Canada ralentis-sait, est-ce qu'il n'y aurait pas un danger que les autres pays

prennent le dessus et qu'en fin de compte, nous soyons à la traîne et nous en subissons les conséquences? J'imagine que c'est le problème; si un pays le fait et l'autre ne le fait pas, ça devient la course folle pour tous les pays.

Comment sommes-nous capables de répondre à ce défi? Je pense que retirer le Canada de la course va nous mettre dans une position précaire. Est-ce qu'il y aurait un moyen de faire les choses pour que tout le monde s'entende et aille vraiment dans le bon sens en ce qui a trait à l'intelligence artificielle, et surtout à la superintelligence artificielle?

• (1120)

Le président: Je vous demande une réponse rapide de 20 secondes ou moins, monsieur Alfou.

[Traduction]

Gabriel Alfou: Je dirais que vous venez de décrire trois préoccupations distinctes. Premièrement, le fait de mettre un terme au développement de la SIA, la superintelligence artificielle, ne nuira pas au Canada. Aucun pays ne pourra contrôler la SIA à lui seul. Tout pays qui développera un système de superintelligence verra sa sécurité nationale menacée.

Le président: Je suis désolé, mais je dois vous arrêter ici.

[Français]

Madame Lapointe, la parole est à vous pour six minutes.

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

Bienvenue aux témoins.

Je dois dire que vos témoignages mettent en évidence les risques associés à l'intelligence artificielle. D'habitude, nous entendons parler de ses bons côtés, mais, ce que vous essayez de souligner, ce sont les côtés plus difficiles.

J'aimerais que, par vos réponses à mes questions, vous nous aidiez à savoir si le Canada est sur la bonne voie en matière de supervision de l'intelligence artificielle. Nous savons bien que nous sommes très avancés sur le sujet, principalement dans la région de Montréal.

Voici ma première question.

Le Canada a lancé l'Institut canadien de la sécurité de l'intelligence artificielle, ou ICSIA, qui a pour mandat de tester et d'évaluer de façon indépendante les systèmes d'intelligence artificielle avancés. Selon vous, quelle est l'importance pour les pays de créer des instituts publics et indépendants comme l'ICSIA?

[Traduction]

Gabriel Alfou: Il est très important de connaître l'état d'avancement de l'intelligence artificielle. Cependant, je pense que nous avons déjà largement dépassé ce stade à bien des égards.

Des experts nous ont déjà mis en garde contre les risques d'extinction de l'humanité que pose l'IA. De nombreux instituts spécialisés en sécurité et sûreté de l'IA ont établi que certaines IA sont déjà capables de nous convaincre, de nous manipuler et parfois même d'échapper à leur confinement. À mon avis, les résultats obtenus par les systèmes existants sont déjà préoccupants. Des experts nous ont également mis en garde contre les systèmes qui seront bientôt disponibles, d'ici 3 à 10 ans.

C'est important, mais j'estime qu'il l'est encore plus de passer à l'étape suivante, qui doit consister à aller au-delà des simples constats et à prendre des mesures concrètes.

[Français]

Linda Lapointe: Vous parlez des risques associés à l'intelligence artificielle. Quand vous parlez de mesurer, qu'est-ce que vous mesureriez, plus précisément? Je veux comprendre ce que vous voulez mesurer.

[Traduction]

Gabriel Alfour: Je pense que c'est une excellente question.

Je dirais que nous avons déjà franchi de nombreuses limites dangereuses. Par exemple, nous avons déjà des systèmes très persuasifs. Si l'on voulait s'assurer que les systèmes ne peuvent pas manipuler les gens... il existe déjà des systèmes qui excellent sur ce plan. C'est la même chose pour le piratage. Nous avons perdu la partie d'avance si nous espérons que les systèmes actuels n'excelleront pas dans le piratage. Il y en a qui en sont déjà là.

Aujourd'hui, nous ne mesurons que leurs progrès, leurs performances surhumaines et l'avance des uns sur les autres. Nous avons déjà franchi plusieurs étapes dangereuses. Nous sommes déjà dans une spirale infernale qui nous rapproche dangereusement d'un point de non-retour. C'est ce dont nous parlons quand nous évoquons les mesures effectuées.

Intervient aussi, et c'est pertinent, la capacité de l'IA à se développer en toute autonomie. À l'heure actuelle, certaines entreprises utilisent de plus en plus l'IA pour développer d'autres IA. Les humains interviennent de moins en moins dans le processus. C'est d'ailleurs l'un des autres aspects que nous essayons de mesurer, soit le nombre réduit d'êtres humains nécessaires au développement de l'IA. Cette mesure est intéressante, car elle permet de déterminer à partir de quel moment le développement de l'IA pourrait tomber dans un cercle vicieux, l'IA évoluant plus rapidement que nous ne pouvons le prévoir. Voilà quels sont, généralement, les types de mesures qui nous intéressent.

[Français]

Linda Lapointe: Ma question renvoie à une autre question que mon collègue vous a posée un peu plus tôt. Croyez-vous que, si nous nous organisons avec tous les autres pays, nous pourrions éviter les risques que vous énumérez depuis tantôt?

[Traduction]

Gabriel Alfour: Personnellement, je le crois. Ces risques sont concentrés dans des systèmes superintelligents. Je pense que le plus difficile est de surveiller et de réglementer les précurseurs de ces systèmes superintelligents, mais si nous le faisons, il y aura beaucoup d'avantages à tirer de l'intelligence artificielle. Je ne dis pas que ce sera est facile, mais c'est très faisable. Nous pouvons le faire scientifiquement, et je pense que nous devrions le faire.

• (1125)

[Français]

Linda Lapointe: Merci.

Monsieur Leahy, avez-vous quelque chose à ajouter au sujet des risques de l'intelligence artificielle et de son encadrement à l'échelle mondiale?

[Traduction]

Connor Leahy: Je suis d'accord avec tout ce que mon collègue a dit. Une bonne façon de voir les choses est de penser que c'est envisageable, mais difficile à réaliser.

Comme toute autre technologie, l'IA présente de nombreux avantages, mais une course effrénée, notamment entre pays, ne va dans l'intérêt de personne. En fin de compte, personne ne sortira gagnant de la mise en œuvre de la superintelligence. Peu importe qui y parviendra, cela n'apportera aucun avantage. Un marché de l'IA bien réglementé, bien compris et bien contrôlé présenterait de nombreux avantages. Il sera difficile d'y parvenir, mais il n'y a pas d'autre choix.

[Français]

Linda Lapointe: Je dois dire que vos propos m'inquiètent beaucoup. Je suis certaine que nous pouvons y arriver si nous avons de bonnes intentions et si tout le monde se met à la table pour essayer d'établir une réglementation.

Si nous avions des règles mondiales pour encadrer le développement de l'intelligence artificielle, comme vous le dites, il y aurait de meilleures retombées pour tout le monde. Il faut cependant agir à l'échelle mondiale. Est-ce que je comprends bien ce que vous dites?

[Traduction]

Gabriel Alfour: Je dirais que oui. Pour atténuer les risques liés à la SIA, la superintelligence artificielle, il faudrait agir à l'échelle mondiale. Si un pays développe la SIA, nous en souffrirons tous. Nous vivons dans un monde très interconnecté. Si quelqu'un parvient à mettre au point un système superintelligent capable de renverser des gouvernements, de jouer à la géopolitique et à la guerre mieux que n'importe quel être humain, nous nous retrouverons dans de beaux draps.

[Français]

Linda Lapointe: Merci.

Le président: Merci, madame Lapointe.

Monsieur Thériault, vous avez la parole pour six minutes.

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

Messieurs, vos présentations sont assez percutantes.

Je vais d'abord vous poser une question un peu plus technique.

Certaines personnes diraient que ce discours est alarmiste, que l'intelligence artificielle n'est pas pour demain, que nous ne sommes pas encore rendus à la superintelligence artificielle et que nous aurons le temps de nous pencher là-dessus au moment où il le faudra.

Pourtant, deux tendances indiquent totalement le contraire de ce que disent ces gens. Vous me direz si vous êtes d'accord là-dessus. Premièrement, la puissance de calcul augmente de façon exponentielle. Deuxièmement, les modèles d'intelligence artificielle que nous avons actuellement semblent développer leur intelligence de façon exponentielle grâce aux données disponibles et à la taille du modèle.

Êtes-vous d'accord sur le fait que nous n'avons pas autant de temps que nous le pensons?

[Traduction]

Gabriel Alfour: Tout à fait.

Il y a une autre façon de voir les choses: il y a 15 ans, personne n'avait prévu où en seraient aujourd'hui les capacités de l'intelligence artificielle. Le bond a été gigantesque. C'est l'une des principales raisons pour lesquelles les experts nous mettent en garde contre les risques d'extinction de l'humanité. Très peu de gens, et presque personne dans le milieu de l'IA, ne se seraient attendus à ce que nous ayons des modèles aussi puissants qu'un transformateur génératif préentraîné comme en disposent plusieurs entreprises.

Pour faire simple, quand j'étais adolescent, personne ne s'attendait à ce que nous ayons un jour des modèles capables de dialoguer avec nous. C'était inconcevable. Les progrès vont de plus en plus vite.

[Français]

Luc Thériault: Les États ont en quelque sorte une fascination pour les moyens qui permettent des gains d'efficacité, et ils semblent beaucoup plus portés vers la course à l'implantation d'applications pour gagner en efficacité. Nous verrons plus tard s'il y a effectivement des gains d'efficacité.

De ce que je comprends de vos interventions, l'intelligence artificielle est un monstre en développement, et il est évident que nous ne pouvons pas laisser aller son développement sans mieux l'encadrer. Maintenant, il s'agit de savoir comment l'encadrer.

Au Canada, le projet de loi C-27, qui est mort au Feuilleton, proposait la création d'un poste de commissaire à l'intelligence artificielle et aux données au sein d'Innovation, Sciences et Développement économique Canada.

Au sein du récent gouvernement Carney, on a nommé un ministre de l'Intelligence artificielle et de l'Innovation numérique. Nous aimerions bien le rencontrer bientôt, mais il ne veut pas nous rencontrer. Par ailleurs, en juin dernier, il a déclaré qu'il mettrait davantage l'accent sur la recherche de moyens d'exploiter les avantages économiques de cette technologie plutôt que sur la réglementation.

Que pensez-vous de cette approche?

• (1130)

[Traduction]

Connor Leahy: Bien des raisons font que l'on met davantage l'accent sur les aspects positifs que sur les aspects négatifs. Disons-le simplement, c'est franchement plus rentable et plus sympathique.

Il faut être conscient que nous faisons face à des crises multiples à l'échelle de la planète. Je connais moins bien le Canada que, par exemple, le Royaume-Uni, où j'habite actuellement, ou encore que mon pays natal, l'Allemagne ou même les États-Unis, mais je sais que nous faisons tous face à de nombreux défis. On est souvent tenté de penser à des solutions technologiques, et il faut dire que la technologie est souvent une solution très efficace. Elle est un moyen très puissant pour aider à résoudre ces problèmes. Je pense que l'IA peut être utile pour résoudre bon nombre de ces problèmes, mais nous devons rester sceptiques.

À l'époque où l'énergie nucléaire a été mise au point, on considérait qu'elle serait la solution à tout. Certains voulaient fabriquer des avions à propulsion nucléaire, par exemple, qui auraient émis leurs radiations en vol. D'autres voulaient recourir aux explosions nucléaires pour l'exploitation minière et je crois que les Russes l'ont essayé, entre autres choses.

Cela ne veut pas dire que l'énergie nucléaire n'est pas une technologie incroyablement utile et puissante. Les centrales nucléaires sont parmi les moyens les plus efficaces de produire de l'électricité, mais la raison pour laquelle elles sont si sécuritaires, si parfaites et utiles, c'est qu'elles sont encadrées par une excellente réglementation. Cela a exigé des décennies de travail acharné. Il a fallu compter sur les inventions, par de nombreux experts, de formes révolutionnaires d'ingénierie de sécurité pour récolter les fruits de ces percées.

Je dirais que nous assistons à la même chose ici. Si nous essayons de récolter les fruits des percées de l'IA sans une dose adéquate d'ingénierie de sécurité, nous verrons se répéter la même chose que ce que nous avons vu avec les médias sociaux. Nous n'assisterons pas à l'équivalent d'un réacteur nucléaire sûr et économiquement productif.

[Français]

Le président: Il ne vous reste que 15 secondes, monsieur Thériault.

Luc Thériault: Je ne pourrai donc pas poser d'autres questions.

Le président: Merci.

[Traduction]

Monsieur Barrett, c'est à vous pour cinq minutes.

Allez-y.

Michael Barrett (Leeds—Grenville—Thousand Islands—Rideau Lakes, PCC): Selon votre expérience des utilisations malveillantes de l'IA, comme les hypertrucages, les cyberattaques et les outils d'exploitation autonome, quelles mesures de protection ou de réglementation devrions-nous envisager pour protéger les Canadiens? Pouvez-vous nous donner des exemples de mise en œuvre réussie à l'échelle d'une entreprise, d'un pays ou d'une région?

Je vais permettre aux deux témoins de répondre à la question, s'ils le souhaitent.

Gabriel Alfour: Je vais commencer.

On recense essentiellement deux catégories de risques. Premièrement, il y a ceux qui viennent de la superintelligence et, plus généralement, des systèmes qu'on ne peut pas contrôler. Pour ces risques, je recommande essentiellement de réglementer le développement des systèmes.

Le problème avec de tels systèmes, c'est qu'une fois qu'ils sont mis au point, il n'est plus possible de les contrôler; de remettre le génie dans la bouteille. Dans ce cas de figure, la réglementation devient très importante. C'est pourquoi nous mettons beaucoup l'accent sur les accords internationaux et ce genre de choses.

Deuxièmement, il y a ce que j'appellerais les risques plus prosaïques des systèmes actuels. Nous n'en sommes pas encore à la superintelligence, ce qui veut dire que le génie n'est pas encore sorti de la bouteille. Dans ce cas, les choses se jouent davantage au niveau des applications, c'est-à-dire sur ce que font les entreprises d'IA dont il faut réglementer la production. Il vaudrait la peine de mettre en place des règlements rigoureux sur les aspects dangereux de l'IA ainsi que de solides régimes de responsabilité afin que, si une entreprise bâtit un système à des fins malveillantes, elle en soit également tenue responsable.

Je ferais la part entre ces deux éléments, l'un étant le volet mise en place d'une réglementation de la superintelligence artificielle et l'autre étant celui des applications des systèmes actuels et des systèmes dans le court terme.

● (1135)

Connor Leahy: Pour revenir sur ce que mon collègue disait, il convient de souligner que, chaque fois qu'on leur en a laissé l'occasion, ces acteurs ont tenté de blâmer les utilisateurs d'avoir mal utilisé leurs outils.

En matière de responsabilité civile générale, la meilleure pratique consiste à attribuer la responsabilité à la partie de la chaîne d'approvisionnement la mieux à même de remédier au préjudice. Il est évident que, par rapport aux utilisateurs, ce sont les grandes entreprises — qui disposent des meilleurs talents techniques au monde, d'une influence considérable sur les plateformes, etc. — qui sont les mieux placées pour remédier au préjudice. Je m'opposerais à la responsabilisation de l'utilisateur comme moyen de remédier à ces risques et je préconiserais plutôt de mettre la responsabilité sur le développeur ou l'employeur.

Michael Barrett: J'aimerais poursuivre sur la question de la réglementation du développement.

Si nous maintenons les progrès ou le rythme du développement, et que nous signons des traités avec de nombreux autres pays, mais que, par exemple, la Russie ou la Chine, ou les deux, ne participent pas à ces traités, ne risque-t-on pas de se retrouver dans une situation où nos adversaires poursuivront la course à l'IA, tandis que nous resterons sur la touche en espérant que les choses ne dégèneront pas? D'après le tableau que vous avez dressé dans votre déclaration liminaire, nous savons que ce sera presque certainement le cas, mais nous n'avons peut-être pas mis au point des outils qui nous permettront de nous défendre contre cela.

Ai-je raison de dire que nous pourrions aussi utiliser des modèles pour nous défendre contre des États délinquants et les modèles qu'ils pourraient élaborer et déployer à notre détriment?

Gabriel Alfour: Il faut essentiellement voir les choses en deux temps: le temps de la superintelligence et le temps qui la précède.

Dans une situation où la superintelligence serait une réalité, tout le monde serait perdant. Si la Russie bâtit des systèmes d'IA superintelligents, tout le monde y perdra; elle ne sera pas en mesure de les contrôler. Même chose pour la Chine, pour les États-Unis et pour n'importe quel autre pays. C'est une situation où la superintelligence régnera et où seuls les systèmes d'IA superintelligents conserveront une quelconque capacité d'action.

Il y a ensuite la situation où la super IA n'a pas encore été développée, caractérisée par une véritable course. Tant que la superintelligence artificielle n'est pas une réalité, il demeure possible d'orienter la recherche et de tirer de nombreux avantages de la mise au point de systèmes plus solides. C'est la raison pour laquelle les accords internationaux sont importants, mais les choses pourraient tourner au vinaigre si l'on ne parvenait pas à conclure des ententes internationales.

Il en va de même pour la biotechnologie et pour tous les types d'armes dangereuses, comme les armes nucléaires. C'est pourquoi nous croyons que les accords internationaux sont essentiels, car sans eux nous nous retrouverons dans un régime de superintelligence et les choses tourneront mal.

Le président: Je vous remercie de votre réponse.

Et merci, monsieur Barrett, pour vos questions.

[Français]

Monsieur Sari, vous avez la parole pour cinq minutes.

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

Merci aux deux témoins qui sont des nôtres aujourd'hui.

Vous nous avez fait part d'éléments très importants concernant l'intelligence artificielle. Je ne dirais peut-être pas que vos témoignages sont alarmants, mais ils sont très factuels. Je ne peux qu'être d'accord avec vous sur le risque. Je ne peux qu'être d'accord avec vous sur l'ampleur de l'impact que pourrait avoir une superintelligence sur la vie quotidienne de chacun et chacune.

Mon opinion diverge toutefois de la vôtre sur un aspect, et j'aimerais bien que nous en discutions. Vous avez utilisé le verbe « freiner ». Affirmez-vous qu'il faut freiner le développement de ces technologies ou qu'il faut freiner leur usage?

Comme vous l'avez bien expliqué, il existe déjà des technologies ou des systèmes de persuasion. Je pense que les Canadiens font déjà usage de telles technologies. Est-ce que vous voulez qu'on en freine l'usage, étant donné que la majorité de ces systèmes sont développés à l'extérieur du Canada, ou est-ce que vous voulez plutôt qu'on en freine le développement sur le sol canadien?

Vous avez fait un parallèle très intéressant avec la course aux armements nucléaires. Toutefois, dans la course aux armements nucléaires, les frontières géographiques ont leur importance, ce qui n'est pas le cas pour les technologies développées à l'aide de systèmes d'entraînement et de réseaux de neurones artificiels, qui peuvent être utilisées partout dans le monde.

J'aimerais bien entendre votre avis là-dessus.

● (1140)

[Traduction]

Connor Leahy: Ce sont d'excellentes questions qui concernent une situation très difficile.

Le caractère universel des technologies numériques s'accompagne de risques inédits en matière de prolifération et de contrôle. En ce qui concerne les armes nucléaires, par exemple, nous avons en quelque sorte la chance que le minerai d'uranium est assez volumineux et que les centrifugeuses sont plutôt difficiles à construire et assez visibles depuis l'espace, quand elles sont bien conçues. Nous bénéficions de certains de ces avantages en ce qui concerne les systèmes d'IA, tels que les centres de données. Dans d'autres cas, comme avec les codes sources libres, de nombreux systèmes présentent d'autres difficultés.

C'est pourquoi nous insistons autant sur la nécessité de réglementer le développement. Un système superintelligent correspondrait probablement à un fichier informatique qu'il serait ensuite impossible de supprimer et qui aurait la capacité d'essaimer. Il est même possible que nous ne sachions pas à quoi nous aurons affaire après sa mise au point. Il est fort probable que nous ne nous rendions même pas compte que le premier système superintelligent existe déjà, c'est-à-dire quand il sera déjà trop tard. À l'heure actuelle, bon nombre des systèmes d'IA actuels ont des capacités dont nous ignorions tout au moment de leur développement. Ce n'est qu'à retardement que nous découvrons que nos systèmes sont capables de choses que nous ne connaissions pas.

Il s'agit d'un nouveau régime. Par le passé, nous n'avons pas vraiment bien réussi dans des situations semblables, par exemple dans le domaine du contrôle à l'exportation des logiciels dont les résultats n'ont pas été très probants ou dont l'application s'est avérée très difficile.

Précisons qu'il n'y a pas lieu, selon nous, de bloquer toutes les applications ni l'utilisation de l'IA. Je suis certain que mon collègue sera d'accord avec moi pour dire que nous aimons beaucoup les applications de l'IA actuellement disponibles sur le marché. Ce que nous voulons, c'est profiter des avantages que procurent les types d'IA actuelles et continuer à tendre vers des applications plus puissantes.

[Français]

Abdelhaq Sari: Le temps est un peu limité et j'aimerais vous poser une question aussi sur la stratégie canadienne, alors je conclurais cette question-ci en disant que j'espère aussi qu'il y aura un accord international. Vous en avez parlé. Toutefois, je ne suis pas très optimiste à cet égard, parce que je vois la course dans le domaine des serveurs quantiques et des installations. Je pense que je suis un peu moins optimiste que je devrais l'être.

En ce qui concerne la stratégie canadienne, nous avons reçu plus de 11 000 commentaires lors de la consultation nationale. Je crois qu'il faut aussi éduquer les citoyens et les citoyennes. Comment peut-on institutionnaliser la participation citoyenne afin qu'elle devienne un pilier permanent dans l'élaboration des politiques de sécurité liées à l'intelligence artificielle?

[Traduction]

Le président: Veuillez répondre en 30 secondes ou moins, s'il vous plaît.

Gabriel Alfour: Nous croyons que la sensibilisation et l'éducation sont extrêmement importantes pour contrôler l'IA. C'est environ la moitié de notre action. L'éducation vient en premier.

Deuxièmement, il faut faire participer les gens aux décisions concernant le déploiement des systèmes. Beaucoup s'opposent au développement de systèmes d'IA superintelligents, et nous croyons qu'il est bon de faire tout notre possible pour les informer de la situation afin qu'ils aient leur mot à dire.

Le président: Merci beaucoup.

[Français]

Merci, monsieur Sari.

Monsieur Thériault, vous avez la parole pour cinq minutes.

Luc Thériault: Merci beaucoup, monsieur le président.

Le chef de la direction du Machine Intelligence Research Institute nous a dit, quand il est venu témoigner la semaine dernière, qu'un arrêt mondial n'était actuellement pas possible politiquement. C'est pourquoi son organisme se concentre sur la préservation de la possibilité d'arrêter l'intelligence artificielle, en créant une espèce d'interrupteur d'arrêt. La solution qu'il propose est de mettre en place l'infrastructure technique, juridique et institutionnelle nécessaire pour restreindre à l'échelle internationale le développement et le déploiement dangereux de l'intelligence artificielle. C'est ce qu'il appelle un interrupteur d'arrêt. Cela conduirait à un arrêt coordonné, à l'échelle internationale, des activités de pointe en matière d'intelligence artificielle à un moment donné dans l'avenir.

Que pensez-vous de cette proposition?

• (1145)

[Traduction]

Connor Leahy: D'une façon générale, nous estimons que ce qu'on appelle les « boutons d'arrêt » sont très intéressants, pas simplement d'un point de vue technique, mais plutôt d'un point de vue sociopolitique. Par exemple, j'ai déjà demandé à quelqu'un qui travaillait pour une grande entreprise de technologie s'il lui serait possible de couper tous ses serveurs à la demande. Il m'a répondu par la négative, qu'il ne savait même pas où se trouvaient tous les serveurs. Absolument personne dans l'entreprise ne le savait ni quels logiciels tournaient sur ces serveurs.

Pour ce qui est de créer le cadre légal approprié, je ne suis malheureusement pas au courant de la proposition précise dont vous avez parlé. C'est sans doute intéressant, mais il me semble que cela irait à l'encontre du développement.

Nous ne pouvons pas nous contenter d'attendre de voir ce que donnera un système superintelligent, parce qu'il sera déjà trop tard. Il est fort probable que, lorsque le premier système superintelligent sera bâti, nous ne le reconnaitrons même pas comme tel avant un certain temps, et cela arrivera alors beaucoup trop tard.

Il est très important — et c'est pourquoi nous insistons tant sur les précurseurs — de veiller à ce que de tels systèmes ne soient jamais mis au point. Pour ce faire, nous devons nous positionner dans la boucle avant que de tels systèmes ne soient bâtis.

[Français]

Luc Thériault: Vous avez parlé des précurseurs à plusieurs reprises. Comment peut-on restreindre les précurseurs?

[Traduction]

Gabriel Alfour: Je vais répondre à cette question.

Il existe de nombreux types de précurseurs. Certains sont matériels, comme les centres de données ou la manière d'intégrer des unités de traitement graphique dans la chaîne d'approvisionnement. D'autres sont logiciels, comme les types d'IA qui ont été développés et les échafaudages qui y ont été ajoutés.

Il est très important de comprendre qu'il s'agit d'une cible mouvante. Au fil du temps, il devient plus facile de construire des systèmes superintelligents, et de plus en plus d'éléments entrent alors dans la catégorie des précurseurs. C'est également la raison pour laquelle nous pensons qu'il y a péril en la demeure et que nous devons nous attaquer à ce problème sans tarder.

Pour le moment, nous pouvons nous contenter, par exemple, d'empêcher les programmes de recherche visant à créer des systèmes superintelligents, mais nous devrions également nous concentrer sur la limitation des modèles à code source ouvert, car une fois qu'ils sont là, il est impossible de les retirer. Il en va de même pour les centres de données. Chaque centre de données devrait être équipé de boutons d'arrêt et de dispositifs d'arrêt d'urgence. Il devrait y avoir des règles claires sur leur utilisation, etc. C'est le type de réglementation que nous devrions avoir sur les précurseurs.

Fondamentalement, nous avons affaire à une cible mouvante. À mesure que la technologie évolue et que la manière dont elle est conçue change, la cible elle-même change. Les précurseurs d'il y a 15 ans étaient très différents de ceux d'aujourd'hui, et il aurait été beaucoup plus simple de s'attaquer au problème il y a 15 ans.

[Français]

Luc Thériault: Merci.

Combien de temps de parole me reste-t-il, monsieur le président?

Le président: Il vous reste 35 secondes, ou peut-être un peu plus.

Luc Thériault: Je vais poser ma question et, si le temps nous manque, vous pourrez y répondre à mon prochain tour de parole.

Un des aspects éthiques qui me préoccupe est le caractère énergivore des centrales de données. Par exemple, à Bombay, il y a deux centrales au charbon qui sont extrêmement polluantes. Leur fermeture était programmée, mais elles vont continuer de fonctionner pour satisfaire ce besoin gigantesque en électricité des centres de données d'Amazon, qui construit des centres de données partout dans le monde pour rivaliser avec les autres grandes entreprises comme Google. Je suis préoccupé par le fait qu'on développe l'intelligence artificielle au bénéfice des pays riches, mais au détriment de l'environnement et de la santé des gens qui vivent dans les pays en développement.

Le président: Veuillez fournir une brève réponse.

Luc Thériault: Ça peut être une amorce de réponse.

[Traduction]

Gabriel Alfour: C'est là une autre facette du développement de l'IA sans réelle préoccupation pour l'être humain qui n'a pas son mot à dire. C'est essentiellement ainsi que nous voyons les choses.

[Français]

Le président: Merci.

[Traduction]

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

Jacob Mantle (York—Durham, PCC): Merci, monsieur le président.

Je remercie nos témoins de leur présence et de leurs précieux témoignages.

Pour faire suite à la remarque de mon collègue du Bloc au sujet de la protection des êtres humains, je vais m'attarder un peu sur les priorités, notamment sur les actuelles priorités du gouvernement en matière d'intelligence artificielle.

Vous savez peut-être que le gouvernement est en train d'examiner sa stratégie en matière d'IA, mais la stratégie pancanadienne dans ce domaine énumère trois priorités. La première est la commercialisation, c'est-à-dire la monétisation de l'IA; la deuxième concerne les normes ou les mesures de protection dans ce domaine; et la troisième touche aux talents et à la recherche.

Pensez-vous que ces priorités sont dans le bon ordre?

• (11:50)

Connor Leahy: Vous ne serez sans doute pas surpris d'apprendre que nous n'envisageons pas les priorités sous le même angle. Nous sommes tous les deux technologues de formation. Nous aimons la technologie. Nous nous sommes lancés dans ce domaine pour que la technologie contribue à un monde meilleur. La technologie est à double usage. Elle est synonyme de pouvoir. Il est très important de l'utiliser correctement, surtout qu'il s'agit d'une technologie sans précédent qui a ce genre de pouvoir.

Pour parvenir à de bons résultats, le plus important est de bien faire les choses et de ne pas répéter les erreurs commises avec les médias sociaux. Ne permettez pas qu'une technologie existe, simplement pour la laisser exister. Faisons en sorte qu'elle soit au service de l'humanité, ce qu'elle ne sera pas par défaut. C'est ce que nous devons faire. C'est sur cela, selon nous, qu'il faut mettre la priorité absolue.

Côté risque, le plus pressant, selon moi, est la superintelligence.

Gabriel Alfour: Je suis plutôt d'accord.

Je pourrais peut-être ajouter une chose. Je crois que c'est en juillet de cette année que les mini-clips sur YouTube ont franchi le cap des 200 milliards de visionnements par jour. Il y a eu tout un article de la direction de YouTube à ce sujet, qui se disait ravie d'avoir atteint les 200 milliards de visionnements quotidiens.

Nous pensons que bien des créateurs d'IA sont dans ce paradigme. Nous savons que les entreprises d'IA regorgent d'employés techniques. Il est extrêmement satisfaisant pour ces gens-là de voir les pertes diminuer et les résultats des indices de référence augmenter. C'est une bonne chose en soi de miser sur la technologie pour posséder davantage à terme. C'est sympathique. C'est formidable à voir et c'est plutôt facile à comprendre.

Je vais me faire l'écho de ce que vient de dire M. Leahy. Je crois que la plus grande priorité dans le développement des technologies devrait consister à s'assurer qu'elles profitent aux gens, à l'humanité.

Jacob Mantle: Si j'ai bien compris ce vous dites, il faudrait inverser les priorités dans la stratégie du gouvernement pour faire passer les normes et la protection avant la commercialisation. C'est cela?

Gabriel Alfour: Je dirais qu'il y a de nombreuses façons d'aider l'humanité. Dans le contexte de la superintelligence, nous pensons qu'il faut passer par la réglementation et la protection, mais dans d'autres cas, la réponse se situe dans la façon de mettre la technologie au service des personnes.

Personnellement, j'attends beaucoup de l'IA dans le domaine de l'enseignement. Je pense que chacun a ses priorités, mais si l'on essaie d'utiliser l'IA à bon escient, je crois qu'elle peut apporter beaucoup. C'est une technologie très puissante.

Jacob Mantle: C'est très bien.

Le gouvernement du Canada a récemment créé une nouvelle fonction ministérielle, celle de ministre de l'IA. Dans l'un de ses premiers discours, après sa nomination, le ministre de l'IA, M. Solomon, a déclaré que le Canada allait moins insister sur les mises en garde et la réglementation pour faire en sorte que l'économie puisse pleinement profiter de l'IA. Que pensez-vous de cette vision?

Connor Leahy: De façon générale, je crois que l'on peut effectivement tirer de nombreux avantages économiques en négligeant l'épanouissement humain. Nous l'avons régulièrement constaté dans l'histoire. Il existe de nombreuses façons de stimuler le marché boursier au détriment des investisseurs. Par exemple, la déréglementation des pyramides de Ponzi est très rentable à court terme, mais arrive un moment où il y a un coût à payer.

Nous assistons au même phénomène. Parle-t-on de la mise en place d'une intendance responsable à long terme et de l'édification d'une bonne société qui utilise efficacement la technologie? Encore une fois, je pense que le bien de l'humanité passe aussi par l'utilisation de l'IA, mais en l'utilisant correctement et pour faire son bien. C'est plus difficile et moins rentable à court terme.

Le président: Merci, monsieur Mantle.

Monsieur Leahy, je vais vous demander d'abaisser un peu votre microphone. Votre voix semble un peu voilée pour le moment. Nous voulons nous assurer que les interprètes comprennent bien ce que vous dites.

Madame Church, c'est à vous pour cinq minutes.

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

Bienvenue à nos deux témoins.

Je vais commencer par M. Alfour. Que savez-vous de l'Institut canadien de la sécurité de l'intelligence artificielle, l'ICSIA, qui a été lancé en novembre 2024? Avez-vous traité avec cet institut?

• (1155)

Gabriel Alfour: Pas vraiment, à part quelques contacts avec des gens qui, je crois, y travaillaient alors.

Leslie Church: L'un ou l'autre d'entre vous pourrait-il nous faire part de ses impressions sur son mandat et sur la manière dont se présente son travail?

Je signale au Comité que le gouvernement du Canada a créé l'ICSIA en partie pour étudier les risques posés par les systèmes avancés d'intelligence artificielle afin de soutenir l'élaboration d'outils et de lignes directrices pour gérer les risques. Il entretient aussi des collaborations à l'échelle internationale pour élaborer des protocoles de sécurité liés à l'IA.

Il n'existe que depuis un an, mais je me demandais si vous aviez des conseils à nous donner sur la portée de son mandat ou sur la façon dont on pourrait, selon vous, améliorer son travail.

Gabriel Alfour: D'après ce que j'ai compris, l'ICSIA est un observatoire. Il ne réglemente rien et n'impose aucune contrainte à qui que ce soit. Je ne saurais dire si le rôle de l'ICSIA consiste spécifiquement à réglementer et à imposer des conditions. Sa création découle d'une décision politique qui échappe à mon domaine de compétence, mais je pense qu'il y a lieu qu'une entité soit investie de ce pouvoir.

Il y a deux ans, des experts nous ont mis en garde contre les risques d'extinction de l'humanité. Or, nous savons que plusieurs grandes entreprises d'IA se livrent à une course effrénée pour développer une superintelligence. Il est temps d'agir. Que cela passe par l'ICSIA, par le ministre de l'IA ou par une autre entité, j'estime que c'est ce qu'il faut faire.

Leslie Church: C'est l'une des composantes de la boîte à outils dont le Parlement aurait besoin pour légiférer, si jamais nous empruntons cette voie. C'est intéressant, parce que je pense que nous sommes conscients du risque, mais comme vous l'avez mentionné, il sera difficile de comprendre et de trouver des moyens de réglementer l'IA sous ses diverses formes.

Monsieur Leahy, avez-vous une idée de la façon dont nous pourrions élargir le mandat ou le travail de l'ICSIA pour nous aider à élaborer un cadre et des protocoles de sécurité?

Connor Leahy: Je ne connais pas personnellement l'ICSIA. Je n'ai parlé avec personne de cette organisation. Je me suis déjà entretenu avec de nombreux collaborateurs de Yoshua Bengio, à Mila. Il est l'un des architectes de l'IA. Ce sont les gens avec qui j'ai le plus de contacts.

En règle générale, il est très important de pouvoir s'appuyer sur des ressources techniques en mesure de prodiguer des conseils impartiaux aux gouvernements. Il s'agit là d'une fonction essentielle. Il est également très important de comprendre que bon nombre de ces institutions, sans que cela soit de leur faute, sont souvent corrompues par des intérêts commerciaux.

Bon nombre des gens qui travaillent pour ces organismes ou institutions reçoivent des offres très alléchantes de telles entreprises, d'autant que leurs noms sont souvent associés à une technologie bénéfique. Cela étant, l'inverse est également vrai et beaucoup de noms sont rattachés à des technologies qui laissent à désirer. Nous ne voulons ni de l'une ni de l'autre. Ce que nous voulons, c'est un équilibre dans la façon de réduire au minimum les risques vraiment inacceptables, et la façon de tirer profit de ceux qu'il n'est pas possible d'atténuer. C'est un exercice d'équilibre difficile, mais il est important d'établir un mandat clair en ce sens.

Leslie Church: Merci. C'est très utile. Je pense qu'il est utile que les instituts canadiens spécialisés en IA, notamment Mila, collaborent également avec l'ICSIA pour assurer un certain équilibre.

Permettez-moi de vous lancer sur une autre piste. Quel type de contre-protocole proposeriez-vous? Vous avez mentionné les risques géopolitiques que nous font courir des acteurs étrangers. J'aimerais savoir si vous avez des conseils à nous donner sur le plan de la cybersécurité ou à propos des protocoles que nous devrions mettre en place pour protéger nos infrastructures et nos systèmes essentiels. Y aurait-il d'autres contre-protocoles que le gouvernement pourrait étudier?

Le président: Veuillez répondre assez rapidement, s'il vous plaît.

Gabriel Alfour: J'ai tendance à penser que les pays devraient maintenant envisager sérieusement de se doter de pare-feu nationaux pour se protéger contre toute action extérieure. Les pare-feu ont souvent été tabous dans le passé. Si un pays doit assurer sa cybersécurité, la sécurité dans son cyberspace, il se doit d'envisager cet outil.

C'est mon opinion personnelle, pas celle de ControlAI. C'est la réponse la plus directe que je puisse donner à votre question.

• (1200)

Le président: Merci, madame Church.

Monsieur Leahy et monsieur Alfour, je tiens à vous remercier d'avoir comparu devant le Comité ce matin.

Nous allons suspendre la séance pendant quelques minutes, le temps de nous préparer pour la deuxième heure.

• (1200)

(Pause)

• (1205)

Le président: Je déclare la séance ouverte.

Je souhaite la bienvenue à notre témoin pour cette deuxième heure, j'ai nommé Carole Piovesan, associée directrice d'INQ Law;

Nous devions avoir un autre témoin, mais il n'a malheureusement pas pu se présenter. Ce sera donc le « *Carole show* » pour la deuxième heure.

Je vous en prie, madame Piovesan, vous avez jusqu'à cinq minutes pour vous adresser au Comité. Allez-y.

Carole Piovesan (associée directrice, INQ Law): Bonjour. Merci, monsieur le président et distingués membres du Comité.

Je m'appelle Carole Piovesan. Je suis associée directrice chez INQ Law, où je conseille les clients sur la protection de la vie privée, la cybersécurité, la gouvernance des données et la gestion du risque en lien avec l'intelligence artificielle.

J'ai eu le privilège de contribuer aux discussions sur les politiques en matière d'intelligence artificielle à l'échelle nationale et internationale, notamment par l'entremise de l'Observatoire OCDE des politiques de l'IA. J'ai déjà comparu devant votre comité et devant le Comité permanent de l'industrie et du développement des compétences de la Chambre des communes. Je suis professeure auxiliaire à la faculté de droit de l'Université de Toronto, où j'enseigne la réglementation de l'intelligence artificielle. De plus, j'ai cosigné l'ouvrage intitulé *Leading Legal Disruption: Artificial Intelligence and a Toolkit for Lawyers and the Law*.

Les opinions que je vais exprimer aujourd'hui sont les miennes et elles ne reflètent pas celles de mon cabinet.

Pour savoir comment nous devrions gouverner l'IA, nous devons revenir aux principes de base et nous demander ce à quoi doit servir l'IA. En 1950, Alan Turing a proposé ce qu'il a appelé le jeu d'imitation: un test destiné à déterminer si les machines pouvaient penser. Il croyait qu'un jour, les machines seraient capables de jouer à des jeux, de se souvenir, d'observer les résultats de leurs propres comportements, d'apprendre des récompenses et des punitions et même d'introduire délibérément des erreurs dans leur travail.

Aujourd'hui, certains des plus grands chercheurs du monde dans le domaine de l'IA sont divisés à propos de la trajectoire sur laquelle l'IA nous entraîne. Des chercheurs canadiens primés comme Yoshua Bengio et Geoffrey Hinton, tous deux pionniers de l'apprentissage profond, préviennent que nous aurons peut-être bientôt des ordinateurs qui dépasseront l'intelligence humaine, avec des répercussions profondes sur la sécurité et le contrôle... ce qui correspond en fait à la menace existentielle dont nous entendons tous parler. D'autres, comme Yann LeCun, un autre pionnier, avancent un argument qui va plus dans le sens de ce qu'on peut entendre par intelligence artificielle, soit un outil qui augmente l'intelligence humaine plutôt qu'il ne la remplace.

La raison pour laquelle nous nous intéressons à l'IA et les résultats que nous obtenons dans cette quête ont une incidence sur notre façon d'envisager sa gouvernance. Si un outil doit servir à amplifier les capacités humaines, nous devons alors en régir l'utilisation. Et si l'IA est un système autonome capable de raisonner de manière indépendante, il faut en réglementer le développement et le déploiement de manière encore plus vigilante. L'approche du Canada doit tenir compte de ces deux aspects.

De par le monde, au moins trois modèles distincts de gouvernance sont envisagés pour l'IA.

Aux États-Unis, l'administration Trump a opté pour une approche de déréglementation qui met l'accent sur la compétitivité plutôt que sur des mesures de protection globales. L'approche fédé-

rale consiste à s'appuyer sur les lois sectorielles existantes appliquées par des agences telles que la FTC, la Federal Trade Commission, et à s'opposer activement aux expériences menées au niveau des États qui souhaitent instaurer des règles plus strictes en matière d'IA.

Le Royaume-Uni et Singapour ont une approche différente, une approche sectorielle beaucoup plus adaptée à la réglementation de l'IA. Le Royaume-Uni, en particulier, s'appuie sur des principes voulant que les organismes de réglementation sectoriels existants interprètent et appliquent des principes transversaux comme la sécurité, la transparence, l'équité et la capacité de contestation dans leurs domaines. Le Royaume-Uni considère que cette approche lui confère la capacité fondamentale de s'adapter au rythme de l'évolution rapide de la technologie, bien que certains développements laissent entrevoir des mesures contraignantes pour les modèles d'IA les plus puissants.

Singapour a certainement adopté une loi beaucoup plus souple, un cadre à adhésion volontaire. On n'y trouve pas de règlement particulier sur l'IA. Cependant, l'approche de Singapour consistant à rechercher un consensus entre le gouvernement, l'industrie et les citoyens, et à recourir à des instruments comme le cadre de gouvernance modèle de l'IA et la trousse d'outils d'essai de la AI Verify Foundation, s'est révélée relativement efficace pour instaurer la confiance et favoriser une approche commune au développement de l'IA. Compte tenu de son investissement national dans la littératie en IA et de son approche consultative et itérative en matière de gouvernance, le Canada aurait tout lieu de s'inspirer du modèle de Singapour.

Le troisième modèle est beaucoup plus prescriptif. On le retrouve dans la loi de l'Union européenne sur l'intelligence artificielle, dont le Comité a déjà entendu parler. Cette loi est de portée beaucoup plus horizontale et elle porte surtout sur le cycle de vie normatif du développement et du déploiement de l'IA tout au long de la chaîne d'approvisionnement.

• (1210)

L'approche du Canada devrait être adaptée à notre contexte. Réglementer des systèmes d'IA de pointe n'équivaut pas à réglementer l'utilisation de Copilot dans un cabinet d'avocats ou d'un bot conversationnel dans un service d'assistance téléphonique. L'approche contextuelle du Royaume-Uni tient compte de cette réalité. Le Canada ressemble davantage au Royaume-Uni et à Singapour qu'aux États-Unis ou à l'Europe. Nous privilégions une réglementation proportionnée qui protège les droits tout en favorisant l'innovation.

Je terminerai sur un appel à l'action en trois points.

Le premier consiste à continuer d'élaborer une stratégie réglementaire pour une IA sûre. Notre institut dédié à la sécurité de l'IA doit fonctionner à plein régime pour démontrer que le Canada prend au sérieux la sécurité de ces systèmes. Nous devons continuer de privilégier des normes itératives et une approche fondée sur des directives sur l'IA en mettant l'accent sur les essais en conditions réelles pour les IA à haut risque. Les tests de performance en laboratoire et les évaluations hors ligne ne montrent que les performances des modèles lors d'essais statiques, et non leur interaction réelle dans des conditions d'utilisation normale.

Deuxièmement, et c'est très important, nous devons améliorer la diversité des représentations et des perspectives dans les politiques, ainsi que dans le processus de développement, d'évaluation et de mise en œuvre. Les perspectives individuelles sont importantes, et elles sont très peu représentées dans l'écosystème de l'IA.

Troisièmement, nous devons procéder à une analyse contextuelle afin de mieux comprendre, secteur par secteur, les lacunes éventuelles de nos lois en matière d'IA ou encore les domaines où l'IA est déjà prise en compte, cela en vue de ne négliger aucun aspect de l'utilisation quotidienne de l'IA dans les entreprises. Nous devons cibler spécifiquement les lois contraignantes et non contraignantes au niveau national et permettre au Canada de tirer parti du capital de confiance dont il jouit dans le monde pour garantir des normes, des certifications et des orientations solides et harmonisées en matière d'IA responsable.

Merci. Je suis prête à répondre aux questions du Comité.

Le président : Merci, madame Piovesan.

Je sais qu'il y a un quatrième, un cinquième et même un sixième volet à vos recommandations, car j'ai parcouru votre déclaration liminaire. Les membres du Comité pourront peut-être vous y ramener dans leurs questions.

Monsieur Barrett, vous avez six minutes. Allez-y, je vous prie.

Michael Barrett : Je me propose de revenir sur l'un des points que vous avez mentionnés dans votre déclaration liminaire. Il s'agit de la réglementation des systèmes d'IA de pointe. Comment le Canada se défendrait-il contre des pays délinquants qui mettraient au point et déploieraient la superintelligence artificielle? Quelle coordination internationale serait nécessaire pour assurer l'efficacité des mesures de protection contre le développement de systèmes d'IA de pointe?

Carole Piovesan : Chaque nation veut protéger ses propres systèmes dans toute la mesure du possible. Depuis 2019, le Canada travaille à l'échelle internationale dans le cadre du Partenariat mondial pour l'intelligence artificielle de l'OCDE, du G8 et du G7, ainsi que d'autres mécanismes internationaux, notamment les instituts internationaux sur la sécurité en matière d'intelligence artificielle. Nous avons œuvré sans relâche pour dégager une approche commune en matière d'IA et de réglementation, ou de protection des IA de pointe dans le monde, conscients que des pays délinquants pourraient mettre au point des systèmes d'IA allant à l'encontre de nos propres principes.

Il ne faut surtout pas oublier que le Canada siège déjà à ces comités internationaux et qu'il joue un rôle clé dans l'établissement des normes et des mécanismes d'application de ces normes. Nous ne pourrions pas y arriver seuls. Nous devons savoir qui sont nos amis et les points sur lesquels nous avons des approches et des valeurs communes, et, par le fait même, quels mécanismes il nous faudra établir pour défendre ces valeurs et approches.

• (1215)

Michael Barrett : Je voulais vous interroger sur la trajectoire à long terme de l'IA superintelligente et sur les scénarios catastrophes qui pourraient se produire si des mesures de protection adéquates n'étaient pas mises en place pour empêcher cette IA de prendre le dessus sur notre espèce.

Tout d'abord, pourriez-vous me dire de combien de temps nous parlons? Selon vous, s'agit-il du long terme, de 10 ans ou de 5 ans? Qu'en pensez-vous?

Carole Piovesan : Je ne suis évidemment pas technicienne de formation, mais j'ai participé à des cercles de discussion où j'ai pu entendre ce que certains techniciens avaient à dire. Ce que j'ai entendu, et je n'ai aucune raison de ne pas le croire, c'est que nous en sommes à quelques décennies, voire moins. Nous ne parlons plus de...

Je me souviens d'avoir suivi un cours de M. Hinton, il y a des années, lors duquel il avait laissé entendre que l'intelligence artificielle générale n'existerait pas avant des centaines d'années. Puis, en 2023, il a changé de point de vue pour dire que nous sommes beaucoup plus près de cette échéance que nous ne l'avions jamais pensé.

Le développement technologique suit un rythme exponentiel. Rien de tout ce que j'ai entendu ne m'amène à en douter; nous sommes beaucoup plus proches des ordinateurs superintelligents que nous ne l'aurions probablement imaginé il y a quelques années.

Michael Barrett : Il est facile de tomber dans le piège dystopique consistant à se demander à quoi cela pourrait ressembler.

Il me reste un peu plus de deux minutes. Qu'entendez-vous dire sur le sujet? De quoi discute-t-on? Quels sont les résultats que nous voulons éviter? Comment pouvons-nous, en tant que Parlement, contribuer à empêcher que cela se produise?

Carole Piovesan : Nous devons nous garder loin des ordinateurs superintelligents qui nous sont supérieurs au point de nous nuire. M. Hinton a parlé d'instiller un instinct maternel à ces ordinateurs superintelligents afin qu'ils se rendent plus empathiques et protègent.

Le Parlement peut jouer un rôle à cet égard en assurant un suivi régulier de l'élaboration des modèles de pointe et en contribuant à leur élaboration et à leur déploiement à l'échelle de la société canadienne, cela dans le but de définir les valeurs à intégrer dans ces systèmes. Le Parlement peut jouer un rôle en veillant à ce que nous identifions les valeurs que nous souhaitons finalement intégrer dans ces systèmes, puis en mettant en place un processus permettant d'intégrer ces mêmes valeurs dans les systèmes afin de parvenir aux résultats visés.

Michael Barrett : Merci beaucoup.

Le président : Merci, monsieur Barrett.

[Français]

Monsieur Sari, vous avez la parole pour six minutes.

[Traduction]

Madame Piovesan, assurez-vous que l'interprétation fonctionne.

[Français]

Abdelhaq Sari : Merci beaucoup, monsieur le président.

Madame Piovesan, je vous remercie de l'information dont vous nous avez fait part, mais aussi de l'éclairage que vous nous avez fourni.

Je vais quand même commencer par faire la même introduction que dans le cas des deux témoins qui vous ont précédée, en cherchant à savoir ce qu'on peut contrôler.

À votre avis, pourra-t-on contrôler le développement de cette superintelligence? Ne devrait-on pas plutôt travailler davantage à contrôler l'usage même de l'intelligence artificielle, étant donné que la majorité des systèmes qui peuvent persuader les Canadiens et les Canadiennes ne sont pas développés au Canada?

• (1220)

[Traduction]

Carole Piovesan: Bonne remarque. Les systèmes ne sont pas nécessairement développés ici.

Tout d'abord, la coordination internationale nous permet de disposer d'un mécanisme qui permet de déterminer quelles devraient être ces valeurs et ces normes. Il existe plusieurs précédents à cet égard, dont certains ont été évoqués par les témoins précédents. Nous nous devons de tirer au clair ce que signifie l'intégration de ces valeurs positives et ce que pourrait vouloir dire le développement d'ordinateurs intelligents compte tenu du rôle que nous estimons devoir leur attribuer dans la société.

Je vais vous donner l'exemple du G7 de 2018, quand notre gouvernement et tous les gouvernements du G7 parlaient beaucoup des valeurs de l'IA, et de la façon dont nous allions promouvoir des normes, des politiques et des pratiques censées protéger ces valeurs. À partir de là, nous avons assisté à un mouvement dans le sens d'un partenariat mondial sur l'intelligence artificielle, avec l'ONU qui a assumé la coordination en matière d'IA dans le but de parvenir à une plus grande harmonisation. Depuis le dernier sommet du G7, l'accent a principalement été mis sur les questions entourant l'adoption de l'IA, car nous en sommes au point où nous pouvons entrevoir les applications concrètes de l'IA et où nous sommes beaucoup plus enclins et enthousiastes, à bien des égards, à adopter ces technologies, ce qui est une bonne chose.

Dans le contexte actuel, tandis que nous commençons à mieux connaître la technologie et à en comprendre les applications possibles, nous nous devons d'être conscients des risques qu'elle comporte, mais nous ne pouvons pas perdre de vue les perspectives qu'elle offre. Nous avons un rôle à remplir, celui d'assurer une utilisation sécuritaire face à un risque qui est réel. Je ne veux pas me retrouver avec un système qui serait soumis aux mêmes types de contrôles pour toutes les utilisations. Nous devons cibler notre action.

[Français]

Abdelhaq Sari: Exactement.

Vous avez dit qu'il fallait être vraiment conscient des risques, et je partage votre avis à ce sujet. Je pense que notre gouvernement est déjà conscient de ces risques, tout comme le sont plusieurs gouvernements de pays membres du G7, comme vous l'avez dit.

Vous avez aussi parlé de l'approche américaine, qui est beaucoup plus concurrentielle, ainsi que de celle du Royaume-Uni, qui est beaucoup plus basée sur l'éthique et la responsabilisation.

Quelle approche pouvez-vous vraiment suggérer au gouvernement canadien?

Comme j'ai une autre question à vous poser, je vous inviterais à être un peu plus brève, s'il vous plaît.

[Traduction]

Carole Piovesan: Je suis davantage inspirée par les approches du Royaume-Uni et de Singapour que par celles de l'Union européenne ou des États-Unis.

[Français]

Abdelhaq Sari: Ne voyez-vous pas quand même que cette approche pourrait ne pas avoir les résultats escomptés si les autres gouvernements ne sont pas alignés d'une manière ou d'une autre? Cette superintelligence artificielle va continuer à se développer.

[Traduction]

Carole Piovesan: Vous avez tout à fait raison. Cette approche n'exclut pas nécessairement le développement de l'IAG, l'intelligence artificielle générale, mais elle la soumet à une démarche plus mûre dans la façon de réglementer ce secteur.

[Français]

Abdelhaq Sari: Ma dernière question concerne le projet de loi C-27, auquel vous étiez quand même favorable.

Quelles leçons tirez-vous de ce projet de loi et quelle serait la meilleure voie à suivre? Nous avons quand même déjà accompli du travail, et nous ne pouvons pas tout jeter. Selon vous, quelles seraient les leçons apprises en ce qui concerne ce projet de loi?

[Traduction]

Carole Piovesan: Il faut conclure du travail sur la partie 3 du projet de loi C-27, soit la Loi sur l'intelligence artificielle et les données, que le cadre général de responsabilité envisagé aurait exigé de se pencher sur le contexte d'utilisation de l'IA, d'en déterminer le niveau d'impact potentiel, puis d'établir un processus de diligence exhaustive en réponse au niveau de risque.

[Français]

Abdelhaq Sari: Je n'ai pas encore lu votre livre, mais je pense que ce serait vraiment une bonne lecture pour le temps des Fêtes.

Pourriez-vous me dire brièvement quelle est votre position quant à la souveraineté numérique au Canada?

[Traduction]

Carole Piovesan: C'est une question complexe à laquelle je vais devoir répondre en peu de temps.

Je comprends et j'apprécie le concept de souveraineté numérique. Bien des choses sont extrêmement importantes à cet égard. Je reconnais également qu'il s'agit d'un investissement à long terme.

• (1225)

[Français]

Abdelhaq Sari: Merci beaucoup.

Carole Piovesan: Merci.

Le président: Merci, monsieur Sari.

Monsieur Thériault, vous avez la parole pour six minutes.

Luc Thériault: Merci beaucoup, monsieur le président.

Bienvenue, madame Piovesan.

En juin 2025, le professeur Bengio, qui est aussi fondateur de Mila, l'Institut québécois d'intelligence artificielle, et conseiller scientifique au sein de cet institut, a lancé une nouvelle organisation de recherche à but non lucratif sur la sécurité de l'intelligence artificielle, nommée LoiZéro, pour prioriser la sécurité à l'égard des impératifs commerciaux.

D'après vous, quels sont les plus grands risques que pose l'intelligence artificielle pour la sécurité?

Vous avez dit tout à l'heure que, d'entrée de jeu, il faut d'abord se demander quels sont les objectifs de l'utilisation de cette technologie. On est beaucoup dans l'appât du gain.

[Traduction]

Carole Piovesan: Nous devons être conscients des risques de l'intelligence artificielle sur le plan de la cybersécurité, et déterminer si notre pays est prêt à se protéger contre ces risques. Nous devons également prendre acte des préoccupations liées aux droits de la personne et au développement social que soulève l'intelligence artificielle, et garder les yeux grands ouverts à l'heure où les entreprises de notre pays commencent à s'intéresser de plus en plus à l'IA.

[Français]

Luc Thériault: Au sujet des droits de la personne, la Déclaration de Montréal pour un développement responsable de l'intelligence artificielle établit ceci au sujet des SIA, soit les systèmes d'intelligence artificielle:

1. Les SIA doivent être conçus et entraînés de sorte à ne pas créer, renforcer ou reproduire des discriminations fondées entre autres sur les différences sociales, sexuelles, ethniques, culturelles et religieuses.
2. Le développement des SIA doit contribuer à éliminer les relations de domination entre les personnes et les groupes fondées sur la différence de pouvoir, de richesses ou de connaissance.

Le gouvernement tient-il suffisamment compte des biais de l'intelligence artificielle dans la mise en œuvre des politiques? Par exemple, l'a-t-il fait dans le cadre de son projet de loi C-27? Comment devrait-il intégrer davantage cette préoccupation?

[Traduction]

Carole Piovesan: Nous nous devons d'établir un dialogue entre les différents régimes juridiques, tels que le régime des droits de la personne en rapport avec l'IA. En d'autres termes, nous devons comprendre le rôle de l'IA au vu des droits de la personne et de notre charte, par exemple. L'utilisation de l'IA pour produire des résultats équitables suscite des inquiétudes légitimes. Nous devons commencer par comprendre ce que signifie la notion d'équité. Quelle est la norme? Comment l'évaluer? Comment démontrer que nous respectons les normes établies? C'est essentiel.

De plus, je tiens à revenir sur un aspect important que j'ai soulevé plus tôt, celui de la diversité des points de vue autour de la table. Il faut absolument être à l'écoute des divers acteurs qui ont des points de vue différents afin de pouvoir donner forme à notre façon d'appréhender les politiques en matière d'IA sur le plan juridique.

[Français]

Luc Thériault: Revenons à la démarche de M. Bengio. Des garde-fous peuvent-ils être mis en place pour se prémunir contre les risques les plus courants de l'intelligence artificielle? Tout à l'heure, vous avez fait une série de recommandations. Pouvez-vous nous en parler davantage ou nous en présenter d'autres?

[Traduction]

Carole Piovesan: Plusieurs recommandations nous invitent à veiller à ce que l'IA soit utilisée de façon sécuritaire. J'ai essayé de regrouper la plupart de ces recommandations dans mon mémoire, simplement pour accélérer les choses.

Nous devons renforcer les normes en place, comme celles de l'Organisation internationale de normalisation. Le Conseil canadien des normes travaille actuellement à la réalisation d'une série de re-

cherches visant à soutenir les normes canadiennes en matière d'IA. L'application de ces normes est importante.

Les normes ISO sont des normes de gouvernance qui définissent les modalités de mise en œuvre d'une IA responsable en lien avec l'utilisation d'un système. Comme nous parlons en fait de l'utilisation de l'IA dans un contexte particulier, nous ne sommes pas face à une question technique, mais plutôt à une question de gouvernance.

En améliorant notre compréhension de l'IA grâce à un programme national de littératie, à des directives plus sectorielles, à une collaboration accrue entre les industries et à l'apport de points de vue plus diversifiés, nous pourrions commencer à élaborer des plans très concrets sur la manière dont nous allons mettre en œuvre ces normes et pratiques exemplaires afin d'atteindre chacun des objectifs fixés.

● (1230)

[Français]

Luc Thériault: Est-ce que...

Le président: Je suis désolé de vous interrompre, monsieur Thériault, mais vos six minutes sont déjà écoulées.

Luc Thériault: Ah, d'accord.

Le président: Ça passe vite.

Monsieur Hardy, vous avez la parole pour cinq minutes.

Gabriel Hardy: Madame Piovesan, merci d'être des nôtres aujourd'hui.

Plusieurs entreprises ont déjà précisé l'objectif de leur recours à l'intelligence artificielle: remplacer les humains pour une bonne partie du travail qu'il y a à faire. Des études démontrent qu'il y a eu une chute de 13 % des embauches de jeunes travailleurs de 22 à 25 ans dans des emplois où ils sont facilement remplaçables par l'intelligence artificielle.

Pourrait-on dire qu'en décidant de recourir à l'intelligence artificielle comme main-d'œuvre, c'est comme si on avait en ligne des gens de calibre postdoctoral qui peuvent travailler 24 heures sur 24, et ce, pour moins que le salaire minimum? Je crois que c'est déjà quelque chose dont on parle.

L'éthique et l'innovation responsable sont au cœur de tout. Pensez-vous que nous devrions légiférer, dans ce cas-ci, pour nous assurer que les entreprises ne commencent pas à utiliser l'intelligence artificielle au lieu d'engager des humains?

[Traduction]

Carole Piovesan: Je ne pense pas qu'il sera possible de légiférer pour empêcher les entreprises de remplacer des ressources humaines par l'IA. Permettez-moi de vous proposer un point de vue différent.

Nous accompagnons régulièrement nos clients dans la mise en œuvre de leurs programmes de gouvernance de l'IA. Quand nous répertorions leurs cas d'utilisation, je leur pose toujours trois questions. Premièrement, quelles sont les tâches qu'ils ne veulent pas accomplir parce qu'elles sont ennuyeuses et banales? Deuxièmement, combien de temps consacrent-ils à chacune de ces tâches? Troisièmement, que feraient-ils s'ils n'avaient pas à accomplir ces tâches? À chaque fois, ils me répondent qu'ils seraient plus proactifs, qu'ils pourraient mieux remplir leur mission et qu'ils apporteraient davantage de valeur à leur organisation.

Voici où je veux en venir.

Tout d'abord, l'IA sera utilisée par nos entreprises, et nous ne pouvons pas ni ne devrions nous y opposer.

Deuxièmement, nous allons devoir réorienter le marché du travail. J'ai des enfants. Je suis parfaitement consciente de ce qui les attend et des vulnérabilités que pourraient présenter certains de leurs choix professionnels. Je comprends que la structure du marché du travail au Canada va évoluer, mais nous devons nous adapter à ces changements. Au lieu de nous y opposer, nous devons soutenir la reconversion et le perfectionnement professionnels. C'est un thème de discussion nationale depuis de nombreuses années.

Le dernier point concerne la culture numérique. Nous devons permettre aux gens de comprendre comment utiliser l'IA et leur donner les moyens de mieux définir et façonner leur propre parcours professionnel, compte tenu de l'impact transformationnel profond de l'intelligence artificielle.

[Français]

Gabriel Hardy: Je comprends ce que vous voulez dire: si quelqu'un n'est pas en train de faire la partie plate de son travail, ça veut dire qu'elle peut être en train de faire quelque chose de super créatif. Cependant, il faut réfléchir au fait que, si tout le monde est pareil et a les mêmes capacités, l'intelligence artificielle pourra exécuter une bonne partie du travail qui est fait au pays. Ne pensez-vous pas que des entreprises vont essayer d'en profiter? Il faudrait alors s'assurer que les gens qui vont perdre leur emploi ne deviendront pas un fardeau pour la société en général, du fait qu'il faudrait payer pour eux, par exemple.

Ne pensez-vous pas qu'on est capable de réglementer les entreprises qui remplaceraient des employés humains par l'intelligence artificielle?

[Traduction]

Carole Piovesan: Je vais répondre à votre question en deux temps.

Premièrement, devrions-nous prendre des mesures spécifiques pour accompagner cette période de transition, car certaines personnes pourraient être à la recherche d'autres débouchés en raison des bouleversements causés par l'IA? Deuxièmement, devrions-nous prendre des mesures spécifiques pour empêcher les entreprises de remplacer des ressources humaines par l'IA?

Sur le deuxième point, je ne pense pas que nous devrions empêcher les entreprises de s'adapter à l'avènement de l'IA. Je ne pense pas que ce soit la bonne approche. Certes, cette question va au-delà de mon domaine d'expertise, et je n'exprime donc ici qu'une opinion personnelle, mais je ne pense pas que ce soit la bonne approche.

Je suis d'accord avec vous sur la première partie de vos remarques. Nous devons surveiller et étudier en permanence comment mieux soutenir nos travailleurs tandis que nous passons par une phase de transformation progressive.

• (1235)

[Français]

Gabriel Hardy: J'ai une autre question, dans un autre ordre d'idées.

Dans le secteur privé, c'est quand on montre l'exemple à suivre que le reste du marché suit. Que pouvons-nous faire ici, au Canada, pour devenir l'exemple à suivre en matière d'intelligence artificielle

et inciter les pays à nous imiter, plutôt que de rester à la traîne des autres?

Pouvons-nous adopter une meilleure approche sur le plan éthique? Pouvons-nous faire en sorte que le développement de l'intelligence artificielle s'en tienne à des utilisations dans des domaines très particuliers?

[Traduction]

Le président: Veuillez répondre très rapidement, s'il vous plaît.

Carole Piovesan: Je vous remercie de cette question.

Nous pouvons montrer l'exemple de deux façons. Nous pouvons d'abord adopter la technologie, mais aussi faire valoir notre image de marque de pays responsable en matière d'application de l'IA sur laquelle nous travaillons depuis très longtemps. Nous devrions la vendre au reste du monde.

Le président: Merci, monsieur Hardy.

Monsieur Saini, vous avez cinq minutes. Allez-y, monsieur.

Gurbux Saini (Fleetwood—Port Kells, Lib.): Merci de votre témoignage, madame Piovesan.

J'ai été très alarmé par les propos des deux témoins qui vous ont précédée, avec leur description de l'ampleur du risque qui pèse sur l'humanité. Que pouvons-nous faire? Les États-Unis, la Chine, l'Inde et les pays du G7 ont les moyens de développer toutes ces technologies qui pourraient constituer un danger pour l'humanité, mais beaucoup d'autres pays n'ont ni les connaissances ni les installations nécessaires pour le faire. Que pourrions-nous faire pour les aider? Comment pourrions-nous protéger l'humanité contre l'utilisation incontrôlée de ce formidable outil?

Carole Piovesan: Je crois qu'il en a déjà été question ici. Je n'ai pas suivi l'entièreté des témoignages des deux témoins précédents, mais je crois qu'ils en ont parlé dans le contexte de la coopération nucléaire. Je pense que oui. Nous allons devoir nous appuyer à nouveau sur ces processus. Nous devons savoir qui sont nos amis, puis établir avec eux les mécanismes qui permettront de mettre en place...

Le président: Pardon, madame Piovesan, je suis désolé.

Votre son est devenu caverneux. Je me demande si votre micro ne s'est pas déconnecté. Si cela ne vous dérange pas, essayez de le rebrancher, car le son dans la salle est soudainement devenu très caverneux.

Monsieur Saini, sachez que j'ai arrêté le chronomètre. Je vais donner à notre témoin l'occasion de répondre à la question dans son intégralité.

Madame, pouvez-vous dire quelques mots pour tester la qualité, s'il vous plaît?

Carole Piovesan: Bien sûr. Est-ce mieux ainsi?

Le président: C'est nettement mieux. Merci.

Je vais vous donner l'occasion de reprendre votre réponse. Une fois que vous aurez terminé, je redémarrerai le chronomètre.

Carole Piovesan: Fantastique.

Comme je le disais, il a en déjà été question ici. Je pense que nous devons déterminer qui sont nos amis et mettre en place les mécanismes nécessaires pour ancrer les bonnes valeurs, les mécanismes de certification, les mécanismes d'évaluation et les garanties de déploiement afin de faire tout notre possible pour empêcher les pays délinquants d'investir dans des ordinateurs superintelligents nuisibles à l'humanité et, au final, de réussir dans leur entreprise. Je pense que c'est la meilleure voie à suivre.

Gurbux Saini: Je crains que nous ne refassions la même erreur qu'avec le nucléaire dont l'utilisation a causé énormément de ravages avant que nous ne nous rendions compte à quel point il était dangereux.

Des organisations, comme les Nations unies, pourraient-elles réglementer ces choses et dire aux pays délinquants qu'assez c'est assez, que nous ne voulons pas poursuivre dans cette voie?

Carole Piovesan: Je ne pense pas qu'il s'agisse de la bonne organisation internationale de gouvernance pour nous protéger contre l'IA. Nous allons devoir investir dans une nouvelle organisation ou renforcer une organisation existante en vue de soutenir le déploiement sécuritaire de l'IA. Nous avons pu le constater en partie grâce au réseau international des instituts de sécurité de l'intelligence artificielle — qui regroupe et structure les divers instituts nationaux de sécurité —, mais je ne crois pas que ce soit le mandat explicite de ces organismes.

• (1240)

Gurbux Saini: Dans le domaine de l'IA, comment, de façon générale, les entreprises canadiennes se comparent-elles à celles du reste du monde?

Carole Piovesan: Nous avons Cohere dans le domaine des GML, les grands modèles de langage, qui est compétitive à l'échelle mondiale. Sinon, d'après ce que je comprends, nos entreprises d'IA sont relativement petites comparées à certaines entreprises américaines. Je crois que c'est la Californie qui compte 32 des 50 plus grandes entreprises d'IA au monde.

Les entreprises d'IA canadiennes ont encore beaucoup de chemin à parcourir pour rattraper les autres et les rivaliser à l'échelle mondiale.

Gurbux Saini: Je crois savoir que des fonds ont été prévus à cette fin dans le dernier budget. Pensez-vous que c'est suffisant ou que nous devrions en faire plus?

Carole Piovesan: Cela sort un peu du cadre de mes compétences.

Quand on songe à tous les rapports sur la bulle de l'IA et aux sommes colossales investies dans les entreprises américaines spécialisées dans ce domaine, nous ferions bien non seulement d'investir davantage dans certaines de nos entreprises et dans des mesures visant à les soutenir, mais aussi de les aider à se développer à l'international et d'exporter notre technologie dans les règles de l'art.

Gurbux Saini: Vous avez évoqué le modèle de Singapour ou du Royaume-Uni que notre pays devrait suivre. Pouvez-vous nous expliquer pourquoi vous préférez ces modèles à ceux d'autres pays avancés dans le monde?

Carole Piovesan: Le modèle britannique, qui repose sur une approche sectorielle fondée sur des principes, n'exclut pas totalement le concept de réglementation, mais fait appel à une application itérative pour mieux repérer les lacunes réglementaires et déterminer où la réglementation peut être mieux ciblée afin, si besoin était, de justifier une loi de portée horizontale. Sinon, il peut aider les orga-

nismes de réglementation à formuler des orientations efficaces par le biais de consultations auprès de l'industrie et d'autres acteurs. À mon avis, ce modèle cadre très bien avec l'approche générale du Canada.

Singapour est un cas intéressant. Le pays a mis en place plusieurs bacs à sable, notamment AI Verify, qui s'est révélé efficace pour tester différents modèles de confiance. Il a également investi massivement dans un programme national de littératie en IA. Nous observons une réaction différente face à l'IA dans des pays comme Singapour, qui font davantage confiance à cette technologie, ce qui se traduit par une adoption plus large et par une plus grande attention portée à l'économie et à la compétitivité.

Le président: Merci, monsieur Saini et madame Piovesan.

[Français]

Monsieur Thériault, vous avez la parole pour cinq minutes.

Luc Thériault: Merci, monsieur le président.

L'intelligence artificielle est vraiment un monstre en développement, et il est évident que nous pouvons l'encadrer davantage. Nous avons parlé tantôt de l'ancien projet de loi C-27, qui est mort au Feuilleton lors du déclenchement des élections. Il proposait notamment de créer un poste de commissaire à l'intelligence artificielle et aux données.

Que pensez-vous de cette recommandation?

[Traduction]

Carole Piovesan: Je pense qu'un tel poste pourrait être très utile, selon le mandat et surtout selon les ressources qui y seront rattachées.

[Français]

Luc Thériault: L'intelligence artificielle est un catalyseur, voire un accélérateur exponentiel du contrôle de l'information, et donc du pouvoir. Comment peut-on bien encadrer cette volonté inhérente dans les motivations des développeurs?

Le gouvernement, par la voix de son ministre, nous a indiqué qu'il mettrait davantage l'accent sur le développement des avancées économiques de la technologie que sur la réglementation.

Que pensez-vous de cette approche? Ne pensez-vous pas que l'on devrait consacrer autant d'énergie à la réglementation de l'intelligence artificielle qu'à la promotion de ses avantages économiques?

• (1245)

[Traduction]

Carole Piovesan: Si je me souviens bien, je pense avoir dit de l'approche du ministre à l'égard de l'IA, qu'elle était « légère, stricte et juste » et qu'elle revenait à dire que nous chercherions à réglementer, mais d'une manière adaptée et ciblée. Je pense que cette vision cadre très bien avec les approches du Royaume-Uni et de Singapour que je propose pour le Canada.

J'ai participé aux consultations liées à la Loi sur l'intelligence artificielle et les données, et l'on s'inquiétait beaucoup à l'époque de l'établissement d'une loi de portée horizontale qui s'appliquerait dans tous les contextes, mais sans consultations préalables suffisantes pour déterminer si une telle mesure s'imposait.

J'estime que l'approche que je propose au Comité nous permettrait d'établir un équilibre entre les avantages économiques et les contrôles dont vous parlez.

[Français]

Le président: Je m'excuse encore une fois, monsieur Thériault, mais nous avons un autre problème de microphone. Je vais arrêter le chronomètre pendant que nous le réglons.

[Traduction]

Madame Piovesan, pourriez-vous débrancher, puis rebrancher votre casque-micro, je vous en prie? Pour une raison ou une autre, le son de votre voix est creux et métallique.

Pouvez-vous faire un autre test?

Carole Piovesan: Est-ce que cela vous convient?

Le président: C'est beaucoup mieux, oui.

Carole Piovesan: D'accord.

Le président: Je me devais de vous le demander pour protéger l'ouïe des interprètes. Il ne faudrait pas leur causer de blessures.

[Français]

Monsieur Thériault, il vous reste deux minutes dix secondes.

Luc Thériault: Merci, monsieur le président.

Madame Piovesan, vous avez déjà expliqué que la réglementation vise à encourager les organisations à changer leur culture, à améliorer l'éducation, à renforcer la littératie numérique dans toute l'organisation et à adopter une approche plus responsable de ce qu'elles construisent. Vous avez également dit qu'il s'agit d'atténuer les dommages réels en aval. Selon vous, il demeure une part d'incertitude quant à certains règlements, mais, si on prend du recul, on voit les mêmes thèmes revenir sans cesse.

Quels sont ces thèmes?

[Traduction]

Carole Piovesan: C'est vrai. En matière d'intelligence artificielle, nos entreprises font surtout de l'analyse des possibilités et des risques. C'est l'un des thèmes dont j'ai parlé.

Nous voulons qu'elles fassent plus et adoptent l'IA, comme cela se fait ailleurs dans le monde. Les entreprises canadiennes sont encore un peu à la traîne en matière d'adoption. Nous voulons qu'elles commencent à utiliser cette technologie et à gagner en confiance et en aisance avec son utilisation. Nous voulons que cela se fasse dans le cadre d'une utilisation responsable, c'est-à-dire selon une approche fondée sur les risques d'utilisation de l'IA dans des conditions données, en veillant à disposer de contrôles de gouvernance appropriés dans les cas éventuels d'IA à haut risque. Idéalement le contrôle serait assuré par chaque secteur et éventuellement par le gouvernement.

Nous avons vu cela dans le document de synthèse du règlement européen sur l'IA, mais aussi selon une approche sectorielle aux risques potentiels dans certaines conditions. Il y a ensuite l'approche fondée sur les risques propres à une entreprise et la manière dont celle-ci pourrait défendre la classification des risques.

[Français]

Luc Thériault: Tout à l'heure, vous avez parlé d'équilibre. Le rôle principal du gouvernement quant à l'innovation responsable et à la conception responsable ne serait-il pas justement de mettre l'accent sur la réglementation, aussi légère soit-elle, pour s'assurer de sa mise en place, avant même de pousser le développement économique à fond de train?

On veut toujours gagner en efficacité, mais, comme vous l'avez dit tantôt, ces gains d'efficacité peuvent créer d'autres types de problèmes, notamment sur le plan de la main-d'œuvre. On n'a pas nécessairement évalué tout ça.

[Traduction]

Le président: Répondez rapidement, s'il vous plaît.

Carole Piovesan: Oui, mais il nous faut une approche ciblée, et c'est pourquoi je m'inspire des modèles du Royaume-Uni et de Singapour qui sont des exemples à suivre pour réglementer l'utilisation de l'IA et définir la façon de nous y prendre.

• (1250)

Le président: Merci.

[Français]

Merci, monsieur Thériault.

Monsieur Hardy, vous avez la parole pour cinq minutes.

Gabriel Hardy: Merci beaucoup, monsieur le président.

Maître Piovesan, j'ai une question pour vous. Si le ministre avait une seule chose à faire, si l'ensemble de son travail des prochaines années reposait sur une seule décision qu'il aurait prise, selon vous, quelle serait cette décision importante qui permettrait l'évolution correcte du développement de l'intelligence artificielle pour les prochaines années?

[Traduction]

Carole Piovesan: Il s'agirait de renforcer l'importance et le mandat de l'institut de sécurité afin de s'assurer qu'il soit correctement établi pour diriger et soutenir le développement et le déploiement sûrs et responsables de l'IA de pointe.

[Français]

Gabriel Hardy: De façon plus claire, quelle est cette sécurité, selon vous? Qu'est-ce qui doit être la priorité en matière de sécurité sur le plan de l'intelligence artificielle? Dire qu'il faut procéder de manière sécuritaire, c'est assez large.

Est-ce qu'il s'agirait de rendre les entreprises plus responsables, de sorte que, si un jour elles développaient des technologies négatives, elles en seraient tenues responsables?

Est-ce que c'est vraiment en amont que nous devons avoir une structure très rigide pour nous assurer que nous ne dépassons pas les bornes?

Quelle est cette sécurité, selon vous?

[Traduction]

Carole Piovesan: La sécurité dépend des normes qui encadrent le développement de la technologie, de la manière dont la technologie évolue. Elle est aussi fonction d'une définition claire des normes et de l'instauration d'une surveillance continue et en temps réel de l'utilisation de la technologie. Intervient aussi la manière dont nous coordonnons et harmonisons nos efforts à l'échelle internationale pour nous assurer que d'autres pays suivent le mouvement. Voilà tous les aspects sur lesquels je mettrais l'accent.

[Français]

Gabriel Hardy: Tantôt, vous avez dit que le Canada est un leader en la matière et qu'il devrait en venir à vendre son expertise au monde. Présentement, en quoi le Canada se différencie-t-il et est-il un leader en matière d'intelligence artificielle? Comment sommes-nous justement capables d'influencer le monde par nos pratiques?

[Traduction]

Carole Piovesan: Ce que je voulais dire par là, c'est que le Canada s'est forgé une image de pays responsable en matière d'IA à la faveur de sa participation au Partenariat mondial sur l'IA et à l'élaboration des normes ISO de la série 42000 relatives à l'IA. Ce faisant, le Canada a joué un rôle important dans la définition de ce qu'est une IA responsable. Je pense que nous devrions poursuivre dans cette voie, non seulement en vendant les technologies canadiennes ailleurs dans le monde, mais aussi en démontrant que nos entreprises ont mis en place certaines mesures de sécurité qui renforcent la confiance dans la technologie canadienne.

[Français]

Gabriel Hardy: Je trouve ça intéressant, parce que les deux témoins que nous avons reçus un petit peu plus tôt nous disaient qu'au bout du compte, ce sont les incitatifs qui poussent les entreprises à développer l'intelligence artificielle de manière très rapide et dans un but pécuniaire. Nous avons un peu fait le lien avec les réseaux sociaux. Nous réalisons qu'ils ont de gros effets négatifs, notamment sur nos jeunes et sur la santé mentale de nos enfants. Selon les témoins, une fois la technologie développée, les entreprises disent souvent que ce n'est pas leur faute si l'utilisateur en fait un mauvais usage.

Vous parlez de structurer les choses et de s'assurer que les entreprises au Canada développent bien l'intelligence artificielle. Si nous les tenons responsables des aspects négatifs de l'intelligence artificielle et de son développement s'il est fait uniquement dans un but pécuniaire, est-ce que vous croyez que ça limiterait un petit peu le développement de technologies par des entreprises qui, au-delà du fait de générer de l'argent, n'aideront finalement pas la société?

[Traduction]

Carole Piovesan: Cela dépend de la façon dont nous déciderons d'obliger les entreprises à rendre des comptes au sujet de choses qu'elles n'envisagent peut-être pas.

Je pense qu'il serait possible de mettre en place des mesures dissuasives et des mesures incitatives sous la forme de garde-fous et de pratiques exemplaires. Cette formule pourrait s'avérer très efficace pour soutenir nos entreprises.

Cela ne veut pas dire que je suis contre les mesures punitives. Là n'est pas la question. C'est simplement que nous devons comprendre ce que nous voulons réglementer et les outils que nous allons utiliser pour créer un règlement et pour le faire respecter. Il faut moduler.

[Français]

Gabriel Hardy: Je suis assez d'accord avec vous. J'aime faire le parallèle avec les réseaux sociaux, ou même la cigarette: pendant longtemps, on disait que c'était bon pour nous, jusqu'à ce qu'on réalise que ce ne l'était pas. À partir de ce moment, des choses ont été mises en place pour s'assurer qu'on n'en faisait pas la promotion. J'ai l'impression que les réseaux sociaux profitent présentement d'un gros vide juridique. Nos enfants deviennent captifs des réseaux sociaux sans qu'il y ait de côtés négatifs pour ces entreprises.

Je pense que, si nous nous inspirions de ça, nous devrions peut-être diriger le développement de l'intelligence artificielle, pour nous assurer que ce n'est pas l'idée de faire du profit à tout prix qui le justifie.

Est-ce que vous êtes d'accord là-dessus?

[Traduction]

Carole Piovesan: Je suis d'accord, et je pense que nous pourrions être beaucoup plus actifs dans certains domaines pour assurer la transparence et la divulgation en matière d'IA, en particulier relativement à l'utilisation de bots conversationnels publics, qui peuvent être source de confusion pour les utilisateurs. Nous pourrions publier des avis publics afin de clarifier la situation.

Encore une fois, nous devons cibler nos actions et comprendre l'utilisation de l'IA à usage général.

● (1255)

[Français]

Le président: Merci, monsieur Hardy.

Madame Lapointe, vous allez partager votre temps de parole avec Mme Church. Vous avez 300 secondes.

Linda Lapointe: Merci beaucoup, monsieur le président.

Je vous remercie beaucoup d'être des nôtres, madame Piovesan. Les propos que vous tenez sur ce que nous devons mettre en évidence sont très intéressants.

La semaine passée, nous avons entendu un témoin, Antoine Guilmoin, qui disait qu'il connaissait davantage la protection des renseignements personnels. Il nous parlait d'un élément auquel vous avez fait allusion un peu plus tôt, à savoir qu'il y a beaucoup de lois et que nous devrions plutôt nous appliquer à trouver les failles dans nos lois avant d'en faire de nouvelles.

Je sais qu'il y avait le projet de loi C-27, qui est mort au Feuilleton, mais quelles seraient vos recommandations pour pallier les failles en question?

[Traduction]

Carole Piovesan: Je suis d'accord avec cela. Il est recommandé de mener une étude plus poussée, peut-être par l'intermédiaire des organismes de réglementation, comme dans le modèle britannique, en transmettant les informations à un organisme plus central afin de déterminer où se situent les lacunes spécifiques dans l'application de la réglementation dans le contexte de l'intelligence artificielle.

[Français]

Linda Lapointe: Merci.

[Traduction]

Leslie Church: Bonjour, madame Piovesan.

Ma question porte sur l'approche fondée sur les principes et les valeurs dont vous avez parlé.

Quelles pratiques exemplaires avez-vous constatées au Royaume-Uni ou à Singapour en matière d'application de la loi et de transparence? Comment pouvons-nous nous assurer que les normes et les principes en place qui sous-tendent le cadre sont pris au sérieux, sont adoptés et sont appliqués adéquatement? De plus, comment pouvons-nous veiller à ce que les systèmes d'IA soient transparents pour les organismes ou les gouvernements qui appliquent les principes?

Carole Piovesan: Les exemples que nous avons observés — principalement à Singapour et, dans une moindre mesure, au Royaume-Uni — sont le fruit d'un processus de consultation et de collaboration entre les acteurs du marché, pour ainsi dire, et l'organisme de réglementation compétent. Ce processus itératif et continu permet de voir comment ces principes sont appliqués et de repérer les lacunes ou les difficultés rencontrées dans leur mise en œuvre.

À mesure que nous prendrons davantage conscience des domaines où les risques se matérialisent, nous constatons qu'une importance accrue est accordée à la réglementation, accompagnée d'une application plus stricte. Il se peut également que les organismes de réglementation aient besoin de renforcer certaines capacités d'application afin de pouvoir exercer leur compétence dans leur secteur ou leur domaine de réglementation particulier et en matière d'IA.

Leslie Church: Compte tenu des sensibilités commerciales qui existent sans nul doute dans ce secteur, comment pouvons-nous nous assurer que les entreprises font preuve d'une transparence totale et respectent les règles?

Carole Piovesan: Je pense que garantir une transparence totale sera un défi, car de nombreuses entreprises d'IA s'appuient sur la notion de secret commercial pour protéger leur propriété intellectuelle. Cela faisait partie intégrante de la version antérieure du règlement européen sur l'IA. Dans la mesure où nous mettons en place un organisme semblable à un commissaire aux données et à l'IA, nous pourrions vouloir renforcer ce rôle au sein de cette instance particulière.

Leslie Church: L'un des groupes que j'ai rencontrés est Jeunesse, J'écoute — davantage dans le contexte des bots conversationnels et de l'accès public à l'IA en ce moment — qui m'a parlé d'une norme de diligence.

Existe-t-il une solution qui nous permettrait, au rythme du développement et de l'utilisation publique de l'IA, de créer des garde-fous et d'appliquer une norme de diligence, particulièrement dans les situations où nous devons composer avec des préjudices éventuellement causés aux enfants?

Carole Piovesan: Je pense que nous allons assister à une évolution de la norme de diligence, qu'elle s'applique aux enfants, au domaine des soins de santé ou au secteur de l'automobile. Nous allons commencer à la voir changer et être plus spécifiquement appliquée à l'IA.

Le président: Il vous reste 50 secondes.

Leslie Church: Diriez-vous que les entreprises qui innovent dans ce domaine sont tenues à un devoir de diligence suffisant en vertu des lois canadiennes?

● (1300)

Carole Piovesan: Tout dépend du contexte. Dans certains secteurs, on trouve des cadres qui fonctionnent très efficacement pour bloquer les investissements irresponsables en IA au niveau de l'utilisation. Que les entreprises expérimentent à l'interne est une chose, mais la façon dont elles l'appliquent au final en est une autre. Je pense que, dans certains contextes, il existe des mesures de protection efficaces. Je pense aussi que de nombreux secteurs sont en train de s'adapter, la crainte étant alors que les mesures de sauvegarde ne soient pas suffisantes.

Prenons les soins de santé, par exemple. Santé Canada a publié le document *Logiciels à titre d'instruments médicaux* pour nous aider à mieux comprendre comment évaluer les risques associés aux instruments médicaux. Ce n'est pas rien. Le développement de l'IA dans les soins de santé est une chose, mais nous devons suivre ce genre d'approche ciblée pour nous assurer que, lorsqu'elle sera déployée auprès des patients, l'IA répondra à certaines normes, et nous devons savoir quelles sont ces normes.

Leslie Church: Merci beaucoup.

Le président: Merci, madame Church.

Madame Piovesan, je tiens à vous remercier au nom du Comité pour votre témoignage. Je vous remercie de votre contribution.

J'ai deux ou trois choses à dire au Comité. Au cas où vous ne l'auriez pas encore vu — car cela a été distribué à tous les membres —, nous avons reçu une note du cabinet du ministre de l'IA nous annonçant qu'il ne viendra pas témoigner dans le cadre de cette étude. La correspondance se trouve dans le classeur numérique. Je rappelle également aux membres du Comité que le commissaire à l'éthique comparaitra lundi prochain.

C'est tout pour aujourd'hui. Merci à tout le monde.

La séance est levée.

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>