



CHAMBRE DES COMMUNES
HOUSE OF COMMONS
CANADA

45^e LÉGISLATURE, 1^{re} SESSION

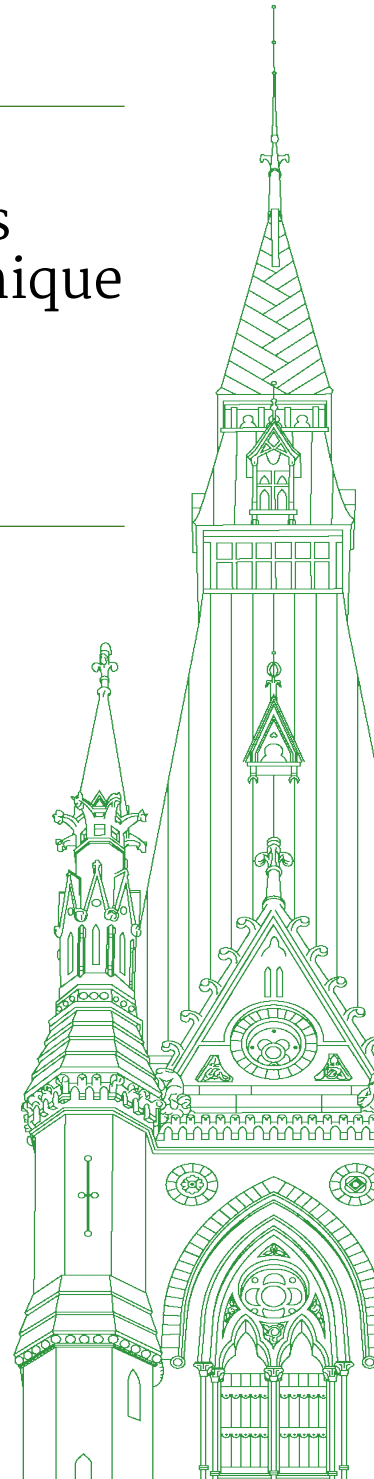
Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique

TÉMOIGNAGES

NUMÉRO 024

Le lundi 26 janvier 2026

Président : John Brassard



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

Le lundi 26 janvier 2026

• (1530)

[Traduction]

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

Je vous souhaite la bienvenue et je vous souhaite également une très bonne année.

[Français]

Bienvenue à la 24^e réunion du Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique de la Chambre de commune.

[Traduction]

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

J'aimerais maintenant accueillir notre témoin pour la première heure de la réunion.

Nous accueillons Wyatt Tessari L'Allié, fondateur et directeur général de Gouvernance et sécurité de l'IA Canada.

Un deuxième témoin, Etienne Brisson, devait être présent aujourd'hui. D'après ce que je comprends, il est aux prises avec des conditions routières difficiles.

Monsieur Hardy, vous avez emprunté le même trajet ce matin.

[Français]

Les conditions routières étaient mauvaises, n'est-ce pas?

[Traduction]

D'accord. S'il arrive à temps pour la deuxième heure de la réunion, je serai heureux de l'inclure parmi les témoins que nous entendrons à ce moment-là.

Monsieur... Je vais vous appeler M. Wyatt, d'accord? Vous avez cinq minutes pour faire une déclaration préliminaire. Vous avez la parole.

[Français]

Wyatt Tessari L'Allié (fondateur et directeur général, Gouvernance et sécurité de l'IA Canada): Monsieur le président, mesdames et messieurs les membres du Comité, je vous remercie de m'avoir fait l'honneur de m'inviter.

Gouvernance et sécurité de l'IA Canada est un organisme à but non lucratif non partisan, ainsi qu'une communauté de personnes d'un bout à l'autre du pays. Notre point de départ est la question suivante: que pouvons-nous faire au Canada et à partir du Canada pour

nous assurer que l'IA avancée est sécuritaire et bénéfique pour tous?

Depuis 2022, nous fournissons au gouvernement fédéral des recommandations de politiques d'intérêt public, telles que nos soumissions sur le projet de loi sur l'IA et les données et nos multiples témoignages aux comités parlementaires.

[Traduction]

Jusqu'à présent, dans le cadre de cette étude, vous avez entendu parler des répercussions auxquelles les Canadiens font déjà face. Même avec les systèmes actuels, des robots conversationnels ont poussé des adolescents au suicide, et les développeurs ne peuvent pas prédire de manière fiable ce que feront les modèles.

Vous avez entendu dire que l'accélération continue des capacités engendre des risques beaucoup plus importants et que des entreprises mondiales comme OpenAI et Google se font concurrence pour bâtir, à court terme, des systèmes d'intelligence artificielle plus intelligents que l'être humain, même si ces entreprises elles-mêmes admettent ne pas savoir les contrôler.

La fusion du cœur d'un réacteur nucléaire est une tragédie, mais le reste du monde continue d'avancer et finit par se rétablir. Toutefois, dans le cas d'une intelligence artificielle plus intelligente que l'être humain, nous n'aurons peut-être pas de seconde chance. Si, à la suite d'un accident ou d'une mauvaise conception, un système interprétait les êtres humains comme un obstacle à la réalisation de l'objectif qui lui a été assigné et qu'il commençait à prendre des mesures contre nous, rien ne garantit que les technologues ou les gouvernements seraient en mesure de reprendre le contrôle de la situation. Cela provoquerait une crise mondiale dont le monde pourrait ne jamais se remettre.

Si vous trouvez cette situation terrifiante, vous n'êtes pas seuls. Il est temps de se demander ce qu'il faut faire. Nous, les Canadiens assis autour de cette table en 2026, que pouvons-nous faire pour faire face à l'avancée exponentielle de l'intelligence artificielle, qui est principalement menée par des entités situées hors de nos frontières?

Si nous le voulons, nous pouvons nous contenter d'intervenir ponctuellement à l'apparition de chaque impact de l'intelligence artificielle et éviter de tenir compte du contexte plus large dans lequel ces impacts s'inscrivent. Nous pouvons tenter de nier ou de minimiser l'importance de ce que les principaux laboratoires sont en train de bâtir, en gaspillant ainsi le peu de temps dont nous disposons pour agir, ou nous pouvons examiner attentivement la direction que prennent les choses et commencer dès maintenant à nous préparer de manière à faire face aux risques actuels, car si le Canada n'est pas prêt à abandonner, il dispose d'un certain nombre de solutions.

En octobre dernier, nous avons publié un livre blanc intitulé *Préparation à la crise de l'IA: un plan pour le Canada*, dans lequel nous formulons quatre recommandations principales.

Tout d'abord, il faut s'adapter pour faire face à la crise liée à l'intelligence artificielle. En effet, la mise au point d'une intelligence artificielle plus intelligente que l'être humain est la plus grande menace pour la sécurité des Canadiens. Pour cette raison, ce dossier mérite de devenir une priorité absolue, car l'intelligence artificielle va bouleverser presque tous les autres dossiers sur lesquels vous travaillez, de la défense nationale à l'emploi, en passant par les soins de santé, l'éducation, l'énergie et l'environnement. Tout comme pour la pandémie de la COVID-19 en 2020, il arrive parfois que le gouvernement doive s'adapter pour faire face à la crise émergente et réévaluer en conséquence l'ordre de priorité des autres dossiers. Compte tenu de sa vaste portée et de ses implications à long terme, la crise de l'intelligence artificielle doit devenir une priorité pour le Cabinet et les mesures prises doivent être coordonnées avec les partis d'opposition et les provinces.

Deuxièmement, il faut lancer la discussion à l'échelle mondiale, car la course à la mise au point d'une intelligence artificielle plus intelligente que l'être humain est un phénomène mondial qu'aucun pays ne peut gérer seul. Aujourd'hui plus que jamais, le monde a besoin d'un chef de file, et le Canada est bien placé pour exercer cette fonction. La meilleure chose à faire est de lancer la discussion et la recherche de solutions à l'échelle mondiale et d'établir les fondements d'un traité sur l'intelligence artificielle que les États-Unis et la Chine pourraient signer lorsque la crise éclatera et qu'ils se rendront compte qu'ils n'ont pas d'autre choix.

Troisièmement, il faut renforcer la résilience du Canada. Si les mesures nationales ne suffisent pas à elles seules à protéger les Canadiens, de nombreuses autres mesures peuvent être prises pour atténuer les répercussions secondaires, par exemple la mise en place de mesures de soutien pour les travailleurs licenciés, l'interdiction des hypertrucages et le renforcement des infrastructures essentielles contre les cyberattaques. En prenant l'initiative à l'échelon national, le Canada sera en meilleure position pour traverser la crise de l'intelligence artificielle et négocier en position de force.

Quatrièmement, il faut lancer un débat national sur l'intelligence artificielle. Les Canadiens méritent d'être informés et consultés au sujet d'une technologie qui va profondément bouleverser leur vie. Nous devons organiser des audiences publiques à l'échelle nationale pour informer et consulter la population sur les décisions de société fondamentales concernant notre avenir face à l'intelligence artificielle.

La semaine dernière, le premier ministre, M. Carney, a donné au Canada un rôle de chef de file sur la scène mondiale. Il s'agit d'une occasion sans précédent de promouvoir l'adoption de mesures de sécurité mondiales en matière d'intelligence artificielle tout en renforçant la résilience au niveau national et en jouant un rôle de premier plan au moment le plus opportun. Les enjeux ne pourraient être plus importants. Le temps presse. Mettons-nous au travail.

• (1535)

Le président: Je vous remercie de votre déclaration préliminaire.

Nous entamons les séries de questions de six minutes avec M. Barrett. Nous entendrons des questions et des réponses.

Pour faciliter le travail des interprètes, je vous demanderais de bien vouloir parler un peu plus lentement lorsque vous répondez aux questions. Je sais qu'ils avaient reçu votre déclaration préliminaire à l'avance, et c'est très bien, mais nous voulons nous assurer que l'interprétation est adéquate.

Monsieur Barrett, vous avez la parole. Vous avez six minutes.

Michael Barrett (Leeds—Grenville—Thousand Islands—Rideau Lakes, PCC): Quels risques fondés sur des données probantes justifient une intervention fédérale urgente du type que vous avez suggérée? Comment éviter que l'élaboration de politiques soit motivée par la crainte plutôt que par la situation sur le terrain? Avez-vous des exemples ou des données probantes à nous présenter?

Wyatt Tessari L'Allié: Oui, certainement.

Nous pouvons actuellement observer une série de signes avant-coureurs, qu'il s'agisse des robots conversationnels qui incitent des adolescents au suicide, des répercussions sur l'emploi des jeunes ou du nombre record d'arnaques et de cyberattaques alimentées par l'intelligence artificielle. Ce sont de graves problèmes auxquels il faut s'attaquer, mais ils ne justifient pas en soi des mesures pangouvernementales. Ce qui importe, c'est la trajectoire de ces problèmes. À l'heure actuelle, nous avons une intelligence artificielle adolescente, c'est-à-dire des systèmes très simples — je devrais plutôt dire relativement simples. Ce que tous les grands laboratoires d'intelligence artificielle sont en train de mettre au point... Si vous consultez leurs sites Web, vous constaterez qu'ils affirment mettre au point des systèmes d'intelligence artificielle plus intelligents que l'être humain. Selon leurs PDG et de nombreux experts, y compris les ingénieurs qui ont quitté ces entreprises pour dénoncer leurs pratiques parce qu'ils ne leur font pas confiance, il sera possible de mettre au point une intelligence artificielle plus intelligente que l'être humain d'ici deux à cinq ans.

Les gouvernements doivent agir de façon proactive dès maintenant, car il se peut que nous n'ayons pas beaucoup de temps pour nous préparer à faire face aux risques bien plus importants qui nous attendent. Étant donné qu'il faudra mettre en place des solutions mondiales et que cela prend une éternité, même si nous avons 20 ans devant nous, nous sommes dans une course contre la montre.

En ce qui concerne l'espoir et la crainte, je suis tout à fait d'accord pour dire que nous ne pouvons pas... Je fais ce que je fais parce que je crois qu'il existe encore des solutions et parce que je pense que des démarches positives peuvent être entreprises. Nous ne pouvons pas laisser la crainte nous paralyser et nous ne pouvons pas affirmer que l'intelligence artificielle est complètement mauvaise, car elle a de nombreuses applications très intéressantes, par exemple dans les domaines des soins de santé et de l'énergie, ainsi que dans d'autres domaines. Je pense qu'il est très important d'éviter d'être trop enthousiaste ou trop pessimiste à propos de l'intelligence artificielle, car il est préférable d'avoir simplement une vision très claire de ce qui nous attend et de la manière dont nous pouvons nous y préparer.

Michael Barrett: Est-il nécessaire de mettre en œuvre plus d'une stratégie pour faire face à certains des effets concrets que nous observons à la suite de l'utilisation de cette nouvelle génération d'intelligence artificielle adolescente? Parmi les effets les plus concrets que nous avons observés, il y a bien sûr les reportages sur les modèles de langage qui conseillent aux personnes vulnérables de se suicider ou de mettre fin à leurs jours. Nous avons également observé une augmentation de la création de contenu sexuellement explicite et d'hypertrucages utilisant des enfants. C'est le pire type d'hypertrucage, mais ce n'est pas le seul. Cela arrive à des personnalités publiques, à toute personne dont la photo est en ligne et, en réalité, à toute personne dont l'apparence peut être décrite dans un modèle de langage. Ce sont quelques-unes des conséquences concrètes que nous observons aujourd'hui. Vous avez aussi évoqué les effets sur le marché du travail pour les jeunes à la recherche d'un premier emploi.

Quelle est la solution? Faut-il prendre une série de mesures? Lorsqu'il s'agit d'hypertrucages, faut-il mettre à jour le droit pénal, afin que les individus soient tenus personnellement responsables de leurs actes pour la création de ce matériel inacceptable qui n'est pas couvert par les lois sur la liberté d'expression et qui va bien au-delà? Puisque cela victimise des gens, faut-il plutôt adopter des lois qui obligent les entreprises de technologie à mettre en place des mesures de protection? Faut-il faire les deux?

• (1540)

Wyatt Tessari L'Allié: Oui, il faut faire les deux.

Je pense qu'il serait utile d'apporter des modifications au Code criminel pour viser précisément les hypertrucages. Je dirais également, comme certains témoins qui m'ont précédé, que les lois actuelles pourraient aller beaucoup plus loin qu'elles ne le font actuellement pour protéger les Canadiens. Si on accordait davantage de ressources aux organismes de réglementation actuels pour qu'ils les appliquent dans le contexte de l'intelligence artificielle, on pourrait peut-être accélérer grandement la mise en place de mesures de protection plus efficaces.

En ce qui concerne les responsabilités liées aux technologies — c'est-à-dire la responsabilité des technologues —, je pense qu'il est très difficile d'empêcher un adolescent de créer un hypertrucage visant une autre personne dans son sous-sol en Russie. Cependant, on peut dire aux intervenants de Google que s'ils souhaitent exercer leurs activités au Canada, ils doivent supprimer les hypertrucages avant une échéance donnée de sorte que même si les images elles-mêmes sont créées, elles ne sont pas diffusées à grande échelle et ne nuisent pas à la réputation des personnes touchées.

Je pense donc qu'il faut certainement prendre ces deux types de mesures.

Michael Barrett: D'accord.

Dans les 10 secondes qu'il me reste, j'aimerais simplement dire que face à tous les risques potentiellement catastrophiques auxquels nous pourrions faire face, nous devrions peut-être d'abord nous attaquer aux défis très concrets et très sérieux comme ceux dont nous venons de parler.

Je vous remercie beaucoup.

Le président: Je vous remercie, monsieur Barrett.

[Français]

Monsieur Sari, vous avez la parole pour six minutes.

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

Bonne année à toutes et à tous.

Je remercie les témoins de leurs présentations.

Avant de poser mes questions sur les façons dont le gouvernement et la Chambre des communes peuvent intervenir, j'aimerais simplement mettre la table en précisant ce dont il s'agit exactement. Quand on parle du monde numérique, d'une manière générale, et de l'intelligence artificielle, d'une manière spécifique, de mon côté, je divise ça en deux volets. Qu'on le veuille ou non, avec la croissance de l'intelligence artificielle, celle-ci devient une infrastructure de base, comme l'électricité, le transport et Internet. Elle est également utilisée dans les processus décisionnels et opérationnels, ainsi que dans les domaines géopolitique et politique. Ça fait augmenter ce qu'on appelle tout simplement la dépendance.

D'une manière générale, quand les citoyens et les citoyennes dépendent de n'importe quelle technologie ou de quoi que ce soit, ça crée une certaine vulnérabilité. En ce qui concerne l'intelligence artificielle, je ne pense pas que nous puissions intervenir sur le plan de la dépendance ou même freiner celle-ci. C'est mon avis, et j'ai quand même travaillé dans le domaine. Cependant, nous pouvons réduire notre vulnérabilité. Nous ne pouvons pas réduire notre dépendance, parce que l'intelligence artificielle est là pour de bon.

Monsieur Tessari L'Allié, selon vous à titre d'expert dans ce domaine, qu'est-ce qu'un gouvernement ou un cadre stratégique gouvernemental peut faire, concrètement? Que devons-nous prioriser pour réduire notre vulnérabilité en réaction à la croissance effrénée de l'intelligence artificielle?

Wyatt Tessari L'Allié: Étant donné la situation dans laquelle nous sommes, où il y a cette accélération, il faut arrêter d'essayer de réagir tout le temps à l'intelligence artificielle d'hier. Ça n'a pas fonctionné. Il faut vraiment être prévoyant. Si ça prend deux, trois ou cinq ans pour mettre en place une mesure gouvernementale, il faut se demander où sera rendue l'intelligence artificielle dans trois ou cinq ans. C'est pourquoi il faut se concentrer sur la préparation plutôt que sur la réglementation.

En ce qui concerne la dépendance, on voit déjà que les travailleurs, par exemple, commencent à oublier comment faire certaines choses, parce que l'intelligence artificielle peut les faire mieux et plus rapidement. Quand on parle de résilience, on ne parle pas juste de l'humain qui oublie comment faire des choses. C'est aussi la société qui perd une certaine résilience, puisque, si cet outil est enlevé ou tombe en panne, tout d'un coup, on ne sait plus comment faire les choses.

Alors, en ce qui concerne la vulnérabilité, ça prend une redondance sociétale. Oui, les outils de l'intelligence artificielle sont très utiles pour améliorer la productivité et aller plus vite, mais il faut maintenir, en même temps, un système d'éducation qui apprend à l'humain comment faire les choses lui-même, afin qu'on ne se retrouve pas dans cette situation de dépendance et de vulnérabilité.

• (1545)

Abdelhaq Sari: Je dis toujours qu'il y a deux aspects. D'abord, il y a la technologie elle-même, qu'on ne peut pas contrôler, parce que, comme vous l'avez dit, quelqu'un peut développer ça dans son sous-sol. Ensuite, il y a les utilisateurs, que nous pouvons quand même éduquer. Vous avez parlé d'éducation, mais j'aimerais bien pousser votre réponse un peu plus loin, sans vouloir vous influencer ou m'immiscer dans votre champ d'expertise.

Y a-t-il un moyen de travailler ou de recommander une certaine souveraineté numérique? Si on n'a pas cette souveraineté numérique, il est difficile de parler des autres technologies que nous ne contrôlons pas vraiment.

Wyatt Tessari L'Allié: Vous avez absolument raison, surtout dans un contexte où nous recevons des menaces des États-Unis. Pour le moment, la majeure partie de notre puissance de calcul est achetée aux États-Unis. Ce sont des centres de données américains. Alors, si jamais on décidait d'éliminer ou de restreindre notre accès à cet outil, tout d'un coup, l'économie canadienne se retrouverait en grande difficulté.

Bref, oui, la souveraineté numérique est importante. Nous devons avoir nos propres centres de données et être capables de faire ce que nous avons à faire chez nous, sans dépendre d'autres pays.

Abdelhaq Sari: Comment pouvons-nous contrôler les Google et les Microsoft de ce monde, alors que nous ne sommes pas encore souverains sur le plan numérique? Il est difficile de s'en sortir. J'ai lu que plusieurs pays, même dans l'Union européenne, essayaient d'avoir cette souveraineté numérique.

En attendant, pourrions-nous exercer un certain contrôle — j'insiste sur le mot « certain » — sur la manière dont ces solutions peuvent être introduites dans notre marché et utilisées par nos concitoyens et nos concitoyennes?

Wyatt Tessari L'Allié: Au risque de me répéter, si on veut ce contrôle, oui, il faut être capable de créer ses propres outils d'intelligence artificielle, d'avoir ses propres centres de données et d'avoir une main-d'œuvre éduquée qui est capable de faire le travail sans dépendre de ça.

Abdelhaq Sari: Je finirai en reprenant ce que j'ai dit tout à l'heure.

À mon avis, comme gouvernement, nous n'avons pas la capacité de contrôler la dépendance des gens envers l'intelligence artificielle et l'évolution de celle-ci. Cependant, comme vous l'avez bien dit, nous avons la capacité d'éduquer nos concitoyens et nos concitoyennes, afin qu'ils soient moins vulnérables, et de faire en sorte que les solutions qui sont utilisées pour tous nos systèmes, que ce soit en matière de santé, de finance ou de transport, soient plus sécuritaires. C'est important, car nos données sont là-dedans, et donc notre histoire et notre avenir aussi.

Merci beaucoup, monsieur Tessari L'Allié.

Le président: Merci, monsieur Sari.

Monsieur Thériault, vous avez la parole pour six minutes.

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

Monsieur Tessari L'Allié, ma première question est peut-être un peu large. Je ne veux pas vous heurter non plus.

Dans votre document, vous proposez de renforcer les cyberdéfenses des infrastructures critiques, de développer les capacités de

défense contre l'intelligence artificielle et les drones, et de préparer les agences de sécurité à la prolifération d'armes biologiques, nucléaires et chimiques, entre autres, facilitée par l'intelligence artificielle générale. C'est une posture défensive — vous en conviendrez avec moi — face à des systèmes qui sont de plus en plus puissants.

On peut bien être d'accord pour se défendre, mais est-ce que cette stratégie n'est pas éventuellement perdante? En d'autres termes, ne faudrait-il pas, de toute façon, empêcher éventuellement le développement de systèmes plus puissants que nos défenses?

Wyatt Tessari L'Allié: Oui, absolument. Il y a l'ère précédant l'intelligence artificielle de niveau surhumain, puis il y a l'ère d'après. À court terme, en ce qui concerne la défense nationale contre les drones, l'intelligence artificielle et tout le reste, oui, il faut augmenter les investissements dans les infrastructures critiques pour les protéger contre les cyberattaques. Il est vrai que, si on atteint un stade où l'intelligence artificielle est capable de nous déjouer sur tous les plans, il n'y a rien que le Canada pourra faire tout seul pour se défendre.

Notre meilleure stratégie, et c'est vraiment la force du Canada dans tout ça, c'est de faire avancer les pourparlers mondiaux, parce que tous les pays sont dans la même situation que nous. Les États-Unis ne pourront pas se défendre contre une intelligence artificielle super intelligente, pas plus que la Chine ou la Russie. Tout le monde a une raison de collaborer à ce sujet, parce que tout le monde est vulnérable à ça.

Bref, essentiellement, je suis d'accord. À court terme, il y a beaucoup de choses que nous pouvons faire pour améliorer notre résilience sur le plan de la sécurité, mais, à moyen terme, la seule option est une solution mondiale qui arrêterait la création de ces systèmes, au moins jusqu'à ce que nous puissions les contrôler.

• (1550)

Luc Thériault: Ce que vous nous dites actuellement, c'est que, selon vos connaissances, on est en train de perdre le contrôle, et qu'à moyen terme, il est fort possible qu'on ne puisse pas contrôler cette technologie.

Wyatt Tessari L'Allié: Oui. Je me réfère aux experts dans le domaine, puisque, pour ma part, je suis plus généraliste et je m'intéresse à l'intersection entre le gouvernement et la technologie. Les experts qui travaillent dans les laboratoires d'intelligence artificielle de pointe, ceux qui sont en train de développer ces systèmes, nous disent qu'ils ne savent pas comment les contrôler et qu'ils ont peur. La seule raison pour laquelle ils accélèrent leur développement, c'est qu'ils pensent que quelqu'un d'autre le ferait à leur place s'ils ne le faisaient pas et qu'ils ont peut-être de meilleures chances de les contrôler. Cette course n'a aucun bon sens.

Pour l'instant, nous avons des systèmes dont les répercussions sont importantes, mais pas bouleversantes. Ce n'est pas la fin du monde. Si ChatGPT fonctionne mal, peut-être qu'un adolescent va mourir ou qu'il y aura une cyberattaque, mais on peut s'en remettre. Cependant, si l'intelligence artificielle atteint un stade où elle peut comprendre ce que nous faisons mieux que nous, agir plus rapidement que nous et nous déjouer sur tous les plans, à ce moment-là, nous serons vulnérables à ses décisions, et nous ne serons plus dans une situation où nos forces de sécurité pourront nous défendre.

Luc Thériault: Ça, c'est sans compter sur ce dont nous parlions tout à l'heure, lorsque les premières questions ont été posées, c'est-à-dire d'une utilisation malfaisante de cette technologie, au-delà des effets négatifs, même délétères, qu'elle peut avoir sur des adolescents, par exemple.

Wyatt Tessari L'Allié: Absolument.

À court terme, ce sont les risques biologiques, par exemple. On peut déjà utiliser l'intelligence artificielle pour créer des nouveaux virus contre lesquels le corps humain aurait du mal à se défendre. Donc, il y a les armes nucléaires et biologiques, et tout le reste. Si l'humain peut concevoir quelque chose, une intelligence artificielle de plus en plus compétente sera aussi capable de concevoir cette chose et de l'utiliser. En effet, même si on ne pense pas qu'on pourrait faire le contrôle à l'intelligence artificielle, le simple risque que celle-ci soit en train de développer une arme de destruction massive est une raison de prendre ça très au sérieux.

Luc Thériault: Êtes-vous optimiste en ce qui concerne les divers gouvernements ou les diverses puissances partout dans le monde, le positionnement géostratégique et géopolitique et les intérêts des superpuissances à devenir encore plus superpuissantes?

Tout le domaine de l'intelligence artificielle est peu encadré actuellement. On est dans un trou noir, contrairement à l'armement nucléaire. Quand on regarde l'armement nucléaire, on peut voir tout ça. Qu'est-ce qui fait qu'on pourrait le contrôler?

Wyatt Tessari L'Allié: Il y a vraiment déjà une course entre pays pour d'autres raisons. Les États-Unis parlent de vouloir dominer en matière d'intelligence artificielle. La Chine veut être chef de file mondial en intelligence artificielle.

Luc Thériault: Quel est le but poursuivi réellement? Quand on parle des États-Unis et de la Chine. Quel but poursuivent-ils?

Wyatt Tessari L'Allié: Pour l'instant, c'est plutôt un but militaire et économique, dans le sens où ils veulent une intelligence artificielle pour pouvoir gagner des guerres et se défendre contre d'autres et pour avoir l'économie la plus forte. Il y a donc des incitatifs très forts à l'échelle nationale pour être premier en intelligence artificielle. Les mêmes incitatifs existent pour les compagnies elles-mêmes; il y a une compétition entre elles. Tout ça nous accélère vers un point de rupture.

Ce qui me rend optimiste, et je dirais plutôt réaliste, c'est que toutes ces compagnies et tous ces pays ont le même problème que nous. Si quelqu'un crée une superintelligence qu'on n'est pas capable de contrôler, tout le monde perd. Que ce soit Donald Trump ou Xi Jinping, ils seront obligés de collaborer à un moment donné, parce que, s'ils ne le font pas, ils vont perdre.

Le président: D'accord. Merci.

Votre temps de parole est écoulé, monsieur Thériault. Vous avez eu plus de six minutes.

Nous allons commencer le deuxième tour de questions.

[Traduction]

Monsieur Cooper, vous avez la parole. Vous avez cinq minutes.

Michael Cooper (St. Albert—Sturgeon River, PCC): Je vous remercie, monsieur le président. Je remercie également le témoin.

Dans votre témoignage, vous avez mentionné les robots conversationnels dotés d'intelligence artificielle, et j'aimerais approfondir un peu ce sujet, en particulier en ce qui concerne les jeunes.

Tout d'abord, je présume que vous conviendrez qu'il s'agit d'un domaine qui doit faire l'objet d'une réglementation. Lors d'une audience d'un sous-comité judiciaire américain qui s'est tenue l'automne dernier, ainsi que dans d'autres rapports, on a présenté des données américaines provenant de Common Sense Media selon lesquelles 72 % des adolescents américains ont utilisé un compagnon doté de l'intelligence artificielle au moins une fois.

Avez-vous des données pour le Canada? Je présume qu'elles seraient similaires.

• (1555)

Wyatt Tessari L'Allié: Oui, je pense qu'elles seraient similaires, mais je n'ai pas vu de chiffres exacts.

Michael Cooper: On vous a interrogé, de manière générale, sur certains des préjudices causés par l'intelligence artificielle, et en particulier par les robots conversationnels. Pourriez-vous nous en dire plus sur certains de ces préjudices chez les jeunes?

Wyatt Tessari L'Allié: Bien sûr.

La question des relations et de la santé mentale est très importante, parce que nous sommes en train d'effectuer cette énorme expérience avec nos enfants en les exposant à ces compagnons dotés de l'intelligence artificielle. Nous ne savons pas comment leur capacité à socialiser et à travailler avec d'autres sera touchée. Nous constatons déjà beaucoup de problèmes de santé mentale. Je suis déçu que l'autre témoin ne soit pas ici, car il a évidemment beaucoup plus d'expérience dans ce domaine.

Les enfants développent une dépendance à l'intelligence artificielle: ils n'arrivent pas à renoncer à ces systèmes qui sont toujours très agréables et obséquieux et nous disent ce que nous voulons entendre. Les enfants se font prendre dans ce cercle vicieux et s'enfoncent dans ces trous sombres.

Cette expérience est en cours, et nous attendons toujours de voir quels en seront les effets à long terme. C'est très préoccupant, car il s'agit d'une période importante du développement des enfants. Si leur santé mentale et leur apprentissage sont perturbés, c'est un problème.

Michael Cooper: Ces systèmes sont entraînés à l'aide de l'ensemble de l'Internet. Est-ce juste? Est-ce exact? Ainsi, ils s'alimentent de tout ce qu'on y trouve: des forums sur le suicide aux sites pornographiques en passant par d'autres contenus préjudiciables. Ces éléments se retrouveront inévitablement — et se retrouvent — dans les conversations que les jeunes entretiennent avec les systèmes d'intelligence artificielle.

Wyatt Tessari L'Allié: Oui, je pense que certaines entreprises s'efforcent d'encadrer les données qui alimentent ces systèmes, mais le problème, c'est que plus le modèle est général, plus il est performant. C'est-à-dire que plus on l'alimente d'une variété d'informations, plus il comprend comment tous ces éléments s'emboîtent. Je ne saurais dire quelles données précises se retrouvent dans les modèles, par exemple, de Gemini ou de ChatGPT, mais le fait qu'ils ont poussé des adolescents au suicide et le fait qu'ils fournissent à l'occasion des instructions sur la façon de fabriquer une bombe chimique laissent supposer qu'ils ont été entraînés à l'aide de contenus très dangereux.

Michael Cooper: Seriez-vous d'accord pour dire que ce qui accentue les risques et réduit la capacité des parents de détecter que leur être cher est exposé à du contenu sexuellement explicite ou préjudiciable, c'est que ces éléments passent souvent sous le radar et que, dans certains cas, il n'y a en fait aucun moyen de savoir que l'on est en train d'échanger avec un système d'intelligence artificielle?

Wyatt Tessari L'Allié: Bon nombre d'enfants préfèrent discuter avec ChatGPT parce qu'ils se sentent plus en sécurité que lorsqu'ils parlent à un parent. Ils ne sont pas au courant des problèmes relatifs à la protection des renseignements personnels, et ils ne savent pas que ces systèmes ne sont pas aussi fiables ou, espérons-le, aussi sages que leurs parents.

Oui, ce problème évolue très rapidement et peut déraiper d'une foule de façons. Nous en constatons déjà les répercussions, et nous nous attendons à ce qu'elles s'accroissent.

Michael Cooper: À l'automne, les sénateurs Hawley et Blumenthal ont présenté un projet de loi bipartite, la GUARD Act, au Sénat américain. Savez-vous de quoi il s'agit?

Wyatt Tessari L'Allié: Plus ou moins.

Michael Cooper: D'accord. En quelques mots, ce projet de loi contient trois objectifs principaux. Il interdirait les compagnons dotés de l'intelligence artificielle aux mineurs, obligerait les robots conversationnels dotés d'intelligence artificielle à divulguer qu'ils ne sont pas humains, et prévoirait de nouvelles sanctions pour les entreprises qui fabriquent des systèmes d'intelligence artificielle destinés aux mineurs et qui sollicitent ou produisent du contenu à caractère sexuel. J'aimerais savoir ce que vous pensez de ces trois objectifs.

Wyatt Tessari L'Allié: Ce sont là, de façon générale, les orientations que vous voulez prendre en matière de protection et de responsabilité, oui.

Michael Cooper: Pour revenir au besoin de transparence, j'aimerais souligner le rapport qu'un membre du groupe de travail du gouvernement sur l'intelligence artificielle a récemment présenté. Il y recommande notamment, en ce qui concerne les produits de l'intelligence artificielle, un étiquetage visible du contenu généré par l'intelligence artificielle, des exigences en matière de transparence des sources, et des mesures relatives aux métadonnées et au tatouage numérique. Seriez-vous également en faveur de ces mesures?

• (1600)

Wyatt Tessari L'Allié: Absolument. Les gens devraient être en mesure de savoir s'ils sont en train d'interagir ou non avec un système d'intelligence artificielle. Ce n'est pas toujours évident en ce moment, et il sera de plus en plus difficile de le savoir, alors oui, étiqueter le contenu généré par l'intelligence artificielle est la moindre des choses.

Le président: Merci, monsieur Cooper.

Madame Church, vous avez cinq minutes. Nous vous écoutons.

Leslie Church (Toronto—St. Paul's, Lib.): Merci, monsieur le président. Monsieur Tessari L'Allié, je vous souhaite la bienvenue.

J'aimerais revenir aux questions de mes collègues. Je suis également très préoccupée par les répercussions de l'intelligence artificielle sur le plan social, surtout chez les enfants. J'ai remarqué dans votre livre blanc que votre approche pour renforcer la résilience du Canada et protéger la sécurité en ligne se concentrait entre autres sur une exigence en matière d'étiquetage du contenu généré par l'in-

telligence artificielle et sur l'interdiction des capacités inacceptables. Je me demande si vous pourriez nous parler un peu de l'étiquetage, de la façon dont vous envisagez son évolution et peut-être de la façon dont il pourrait nous être utile.

Wyatt Tessari L'Allié: Oui. L'étiquetage du contenu généré par l'intelligence artificielle peut se faire de différentes façons. Sur le plan technique, nous n'avons pas de préférence. Cela dit, on pourrait certainement développer une norme mondiale. Par exemple, le protocole Internet qui a contribué à la création de l'Internet est une norme mondiale. Ce serait la même chose avec l'étiquetage du contenu généré par l'intelligence artificielle. Certains ont proposé, à titre d'exemple, de veiller à ce que l'Unicode — le langage qui permet de représenter le texte — soit étiqueté comme ayant été généré par l'intelligence artificielle. Ce faisant, un ordinateur saurait automatiquement si une lettre a été écrite à l'aide de l'intelligence artificielle ou non. De nombreuses solutions de ce genre pourraient se concrétiser.

Au bout du compte, il faudra probablement que le gouvernement force ces entreprises à adopter cette mesure, car l'étiquetage ne sera pas chose facile. Par contre, même si un texte est étiqueté, quelqu'un pourra toujours en prendre une photo et la copier sur un autre ordinateur, faisant soudainement disparaître l'étiquette.

C'est le genre de choses où on ne peut pas... L'étiquetage ne mettra pas fin à ces dérapages, mais il peut faciliter la prise de mesures lorsque l'on remarque que quelqu'un utilise un système d'intelligence artificielle qui n'est pas désigné comme tel. Cette solution ne permet pas de régler tous les problèmes, mais elle constitue un pas dans la bonne direction en incitant les gens à utiliser l'intelligence artificielle de la bonne façon.

Leslie Church: Pouvez-vous me dire ce à quoi cela ressemblerait pour moi, en tant qu'utilisatrice? J'utilise souvent Internet et les médias sociaux; qu'est-ce que je verrais en tant que consommatrice? Comment appliquerait-on cette étiquette?

Wyatt Tessari L'Allié: On pourrait le faire d'une foule de façons.

On pourrait recourir à la couleur de la police. En fonction de sa couleur, on saurait si un texte aurait été généré par l'intelligence artificielle. Une solution semblable pourrait être utilisée pour la voix à la radio. Un bruit de fond précis nous indiquerait qu'il s'agit de l'intelligence artificielle.

Il y a mille moyens techniques de s'y prendre. Il faudrait toutefois que la méthode choisie permette à un utilisateur de savoir instinctivement, dès qu'il voit le contenu, que ce contenu a été généré par l'intelligence artificielle. Il faudrait aussi que ce soit une norme mondiale.

Il y a le logo de la petite étoile à quatre branches. On le voit à différents endroits. C'est un pas dans la bonne direction, mais ce problème très vaste et très complexe n'en est qu'à ses débuts. Les gouvernements du monde entier devront donner des directives claires à ce sujet. Ils devront préciser les exigences et les attentes envers les entreprises qui développent l'intelligence artificielle.

Leslie Church: Lorsque vous dites qu'il faut interdire d'autres capacités inacceptables — je reprends ici le libellé de votre livre blanc —, de quels types de capacités parlez-vous? Pouvez-vous nous en dire un peu plus à ce sujet?

Il va sans dire que nous sommes très inquiets en ce moment. Nous avons présenté une mesure législative relevant du droit pénal pour nous attaquer aux hypertrucages, mais y a-t-il d'autres capacités sur lesquelles nous devrions nous concentrer?

Wyatt Tessari L'Allié: Oui.

Pour reprendre ce que dit la loi sur l'intelligence artificielle de l'Union européenne, les systèmes à risque élevé et moyen dans la catégorie des capacités inacceptables sont, par exemple, les systèmes qui refusent d'être désactivés, qui leurrent l'utilisateur et qui se transforment d'eux-mêmes sans avertissement. Je vous donne un exemple. Si l'on confie une tâche à un système d'intelligence artificielle et qu'il détermine qu'il doit se transformer pour accomplir cette tâche, on se retrouve devant une capacité qui devrait être désignée comme étant inacceptable. S'il pouvait se transformer de lui-même, le système ne ressemblerait plus à ce qu'il était au départ, et le profil de risque changerait radicalement.

Il y a aussi l'autoreproduction. Si on demande à un système de faire quelque chose et qu'il détermine qu'il doit faire des copies de lui-même sur différents serveurs pour que, si le premier serveur qu'il utilise tombait en panne, il puisse continuer à fonctionner, c'est un problème, car, soudainement, le modèle que l'on utilise ne se trouve plus seulement sur notre ordinateur; il se trouve également sur 10 autres ordinateurs auxquels on n'a pas nécessairement accès.

Nos recommandations au sujet de loi sur l'intelligence artificielle et les données contiennent des détails sur les capacités les plus préoccupantes. L'automodification spontanée et l'acquisition de ressources — lorsque le modèle se met à voler des éléments pour accomplir sa tâche — sont autant de comportements que nous commençons à observer en situation de tests, et, si nous les autorisons, nous nous mettons alors dans une situation très vulnérable.

Leslie Church: Je m'intéresse également beaucoup à l'imposition d'une obligation de diligence à ces plateformes. C'est un concept juridique qui, à mon avis, a sa place ici. Il faut mettre en place des garde-fous, en particulier en ce qui a trait aux contenus auxquels les jeunes et les enfants sont exposés, et s'assurer que la responsabilité est clairement établie lorsque les modèles d'intelligence artificielle présentent des renseignements inexacts ou très préjudiciables — comme dans les situations que vous avez décrites —, surtout aux enfants.

• (1605)

Wyatt Tessari L'Allié: La responsabilité doit incomber à ceux qui comprennent le mieux le fonctionnement de ces systèmes, là où se trouvent les meilleurs talents: chez Google, Meta et autres entreprises. Ces entreprises devraient être responsables du comportement de leurs modèles, même après leur déploiement, car on ne peut pas demander à un adolescent ou à quelqu'un qui ne connaît rien à l'intelligence artificielle de savoir quels sont ses dangers et comment l'utiliser.

Leslie Church: Merci.

[Français]

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

Luc Thériault: Merci, monsieur le président.

Monsieur Tessari L'Allié, dans votre document, vous écrivez que, fondamentalement, la gouvernance de l'intelligence artificielle générale est un défi de coordination humaine et de compétence

technique. Le PDG d'Anthropic a récemment dit que nous comprenons peut-être 3 % du fonctionnement interne de ces systèmes.

Si nous ne comprenons pas comment ça fonctionne, pourquoi ne pas nous concentrer sur une négociation pour une interdiction mondiale?

Wyatt Tessari L'Allié: Si nous voyons la crise arriver et que nous ne savons pas comment l'éviter sur le plan technique, notre seule option est d'essayer de nous assurer que la technologie ne sera pas créée en premier lieu. Par exemple, si plusieurs compagnies du centre-ville de Montréal ou de Toronto étaient en train de créer et de développer une technologie nucléaire qui met tout le monde autour d'eux à risque, nous ne serions pas en train de nous demander comment équilibrer la sécurité et les autres aspects; nous enverrions la gendarmerie pour les arrêter.

Dans ce contexte, si nous avons des compagnies en train de créer une technologie qui va mettre tout le monde à risque et qui, elles-mêmes, admettent qu'elles ne savent pas comment contrôler les systèmes qu'elles vont créer, il va falloir que le gouvernement se dise que c'est sa responsabilité de protéger les citoyens et de mettre en pause ces développements, au moins jusqu'à ce que nous sachions comment ces systèmes peuvent être mis en place de manière sécuritaire.

Le rôle principal du gouvernement est donc de protéger les citoyens. Si les technologistes sont en train de faire quelque chose de trop dangereux, le rôle du gouvernement est de les arrêter. Pour ça, nous avons besoin de la participation de tous les pays au monde, parce que, si quelqu'un, quelque part dans le monde, crée ces systèmes, nous sommes en grand danger.

Luc Thériault: Comme vous l'avez dit dans votre document, une solution responsable pour l'avenir de l'intelligence artificielle passe par un accord mondial. Or, cet accord mondial pourrait rapidement devenir une interdiction mondiale. Est-ce exact?

Wyatt Tessari L'Allié: Oui. Il pourrait au moins offrir une pause.

Je suis d'ailleurs très conscient des dynamiques géopolitiques actuelles. Cependant, comme je l'ai dit auparavant, les États-Unis et la Chine sont dans le même bateau que nous: ils ne veulent pas créer de systèmes qui vont leur faire perdre du pouvoir, mettre à risque les citoyens et créer du chômage à grande échelle.

Même s'ils n'ont pas envie de travailler ensemble, ils y seront donc obligés. Pour nous, qui jouons un rôle de pouvoir au sein du groupe, la meilleure action que nous pouvons mener, c'est celle d'essayer de poursuivre les pourparlers et de mettre en place les traités pour que, quand le monde sera prêt à les signer, tout soit prêt pour le faire.

Luc Thériault: Il me reste combien de temps, monsieur le président?

Le président: Il vous reste deux minutes.

Luc Thériault: Monsieur Tessari L'Allié, dans votre document, il est indiqué que les laboratoires pourraient réussir en un an à trois ans à mettre en place une solution mondiale. Or, il nous reste peut-être moins de 18 mois avant que ça arrive.

Compte tenu de ce que nous venons de dire, pouvez-vous donner quelques précisions à ce sujet?

Wyatt Tessari L'Allié: L'avenir est très difficile à prédire. Si on est chanceux, on aura de 10 à 20 ans pour se préparer à ça. On pourra prendre notre temps pour le reste.

Si on est malchanceux, un laboratoire de pointe pourrait dévoiler demain un nouveau modèle qui serait essentiellement la superintelligence. Ça n'arrivera probablement pas, mais peut-être que si. C'est pour ça qu'un gouvernement responsable n'aura pas le choix de se préparer aux scénarios les plus proches.

Si on est chanceux, on aura plus de temps pour améliorer nos solutions. Cependant, il faut commencer. Il faut nous dire qu'on n'a pas le temps. Il faut vraiment agir maintenant, mettre en place les bases, soit en prenant une pause, soit en appliquant d'autres solutions pour s'assurer de ne pas se retrouver dans un monde où on perdra le contrôle.

Ensuite, en ce qui a trait au nombre d'années dont on parle, les PDG d'Anthropic, d'OpenAI et de Google DeepMind parlent de 2 à 5 ans. D'autres experts parlent de un à 3 ans. Nous ne le savons pas, mais, si ce sont des gens intelligents qui suivent de très près ce qui se passe et qui nous disent qu'il faut être prêt — j'espère qu'ils ont tort et j'espère fortement que nous aurons 20 ans pour y arriver — par mesure de prudence, nous devons nous préparer aux scénarios à très court terme. Ensuite, nous prendrons le temps.

• (1610)

Luc Thériault: Nous devons agir par mesure de précaution, finalement.

Wyatt Tessari L'Allié: Voilà, c'est ça.

Le président: Merci. Il vous reste 20 secondes.

Luc Thériault: Je pensais qu'il me restait une minute. Je me suis dit que j'avais dépassé la minute. Très bien, je la reprendrai tantôt.

Le président: Il en reste 14.

Monsieur Hardy, je vous cède la parole pour cinq minutes.

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

Messieurs les témoins, bonjour et bienvenue au Comité.

Souvent, quand on parle d'intelligence artificielle ici, on l'imagine un peu comme dans le film *The Terminator*. On voit la fin du monde arriver. J'ai l'impression que, pour certaines personnes qui suivent nos travaux, ou même pour les décideurs, la situation est improbable et on ne devrait pas trop s'inquiéter. On a dit tout à l'heure qu'on pourrait n'avoir qu'un an et demi pour se préparer, mais peut-être 20 ans. On ne le sait pas.

J'aimerais revenir à un aspect très terre-à-terre, à ce qui se passe au quotidien. Nous en faisons l'expérience. Présentement, l'inflation est très élevée au Canada. Cinquante pour cent des Canadiens sont à 200 \$ près de l'insolvabilité. Il y a 2,2 millions de personnes qui font la queue pour recevoir de la nourriture chaque jour.

Par ailleurs, il y a des entreprises qui voient le potentiel de l'intelligence artificielle pour réduire leur personnel, étant donné que l'intelligence artificielle peut travailler 24 heures par jour, 7 jours par semaine. Selon vous, devrions-nous nous pencher sur cette question dans une perspective étatique? Le gouvernement ne devrait-il pas se dire qu'il doit faire attention étant donné une telle éventualité?

Si on a recours à l'intelligence artificielle au lieu d'engager un Canadien, il y aura des coûts, parce que tous ceux qui perdront leur

emploi finiront inévitablement par avoir besoin des banques alimentaires, qui sont déjà débordées.

Wyatt Tessari L'Allié: Déjà, il y a un an ou deux, Spotify, par exemple, a dit à ses employés de n'embaucher personne avant d'être sûrs de ne pas pouvoir utiliser l'intelligence artificielle.

Des entreprises ont donc déjà donné des instructions très claires à cet égard, parce qu'elles savent très bien que, pour être compétitives, elles vont devoir réduire leurs dépenses. Or, les salaires constituent la dépense la plus importante. Si nous nous retrouvons dans un monde aux prises avec un très haut taux de chômage, la réduction des rentrées d'impôt sur le revenu aura un énorme impact sur le gouvernement.

On remarque aujourd'hui que les répercussions sur l'emploi sont centrées sur certains secteurs tels que ceux qui touchent les informaticiens et les créatifs. C'est un phénomène qui se répandra à l'échelle de l'économie. On ne sait pas si ça arrivera rapidement ou lentement. On verra bien.

Toutefois, si on n'a pas de plan en place pour appuyer ces gens qui se retrouveront sans emploi et veiller à ce que le gouvernement ait assez de revenus pour maintenir des activités — si celui-ci dépend autant des rentrées d'impôt sur le revenu —, ce que nous observons aujourd'hui est vraiment un signe précurseur de ce qui arrivera.

Gabriel Hardy: D'accord.

Je passerais maintenant à un autre ordre d'idées.

Comme vous le savez, le Canada possède le cinquième de l'eau potable mondiale. De plus, je lisais dernièrement qu'un courriel généré par l'intelligence artificielle prend à peu près 500 millilitres d'eau. Si on fait une image, c'est encore plus. Une recherche sur ChatGPT, selon mes notes, consomme dix fois plus d'électricité qu'une recherche faite sur Google.

Est-ce qu'on ne devrait pas, justement, commencer à légiférer sur quelque chose d'aussi tangible que l'eau potable? On s'entend que, si l'intelligence artificielle commence à prendre de plus en plus d'eau potable, c'est d'abord, peut-être, l'humain qui va en subir les conséquences, et ensuite, l'agriculture. J'imagine qu'on est déjà dans une certaine interaction avec les centres de données, puisqu'on en a besoin de plus en plus en raison de l'intelligence artificielle. Ces centres consomment de l'eau et sont en train de nuire, carrément, à la vie humaine.

Wyatt Tessari L'Allié: Oui. Ce qui est très difficile avec les chiffres sur la consommation d'eau et d'électricité, c'est qu'on ne sait pas exactement ce qu'ils sont, parce que les compagnies ne le disent pas. Certainement, dans certaines municipalités où il y a des restrictions d'eau, les centres de données ont un énorme impact sur celles-ci. J'ai entendu d'autres chiffres indiquant que l'impact est presque le même pour un centre de golf, par exemple. C'est pour dire que, effectivement, il y a déjà des répercussions à cet égard. La première étape, c'est vraiment d'obliger les compagnies à nous dire quelle est leur consommation d'eau.

Gabriel Hardy: On s'entend qu'on a besoin de plus en plus des centres de données. Or, plus on utilise l'intelligence artificielle, plus on a besoin de centres de données. Plus on a besoin de centres de données, plus on a besoin d'eau. Si on fait la comparaison avec le golf, ce n'est pas tous les gens qui ont un téléphone qui jouent au golf sur le terrain. Si on fait la comparaison avec l'utilisation de l'intelligence artificielle, chaque personne a cette fonctionnalité sur son téléphone. Si on se dit qu'une petite recherche coûte une bouteille d'eau de 500 millilitres, on peut imaginer qu'on en consomme pas mal chaque jour.

• (1615)

Wyatt Tessari L'Allié: Oui. En ce qui concerne les chiffres sur la consommation mondiale, les États-Unis parlent d'augmenter la consommation d'électricité de 20 % à l'échelle nationale, simplement pour l'intelligence artificielle, au cours des 10 prochaines années. Ça représente des milliards et des milliards de dollars d'investissement. On rouvre des centrales nucléaires pour pouvoir fournir l'électricité pour ça.

Effectivement, c'est quand même à une échelle très importante. Là aussi, je dirais que la première étape, c'est vraiment de demander aux compagnies de nous dire exactement quelle est leur consommation parce que, pour l'instant, il y a beaucoup de fausses nouvelles par rapport à ça. On sait que la consommation existe, mais on ne sait pas exactement ce qu'elle en est.

Gabriel Hardy: Je n'aurai pas le temps de poser ma question.

Le président: Merci, monsieur Hardy.

[Traduction]

Monsieur Saini, vous avez cinq minutes. Vous avez la parole.

Gurbux Saini (Fleetwood—Port Kells, Lib.): Monsieur Tessari L'Allié, je vous remercie de votre présence.

Selon certains, la réglementation de l'intelligence artificielle limitera l'innovation et la compétitivité, tandis que d'autres disent que la certitude réglementaire est essentielle à une mise à l'échelle responsable. D'après votre expérience, une gouvernance efficace de l'intelligence artificielle finit-elle par favoriser ou par limiter l'innovation?

Wyatt Tessari L'Allié: C'est une question difficile, parce que la technologie évolue si rapidement qu'une bonne réglementation en matière d'intelligence artificielle aujourd'hui pourrait être désuète d'ici le mois prochain.

Pour l'heure, dans notre livre blanc, nous ne réclamons pas ouvertement l'adoption de nombreux nouveaux règlements, pour la simple raison qu'il existe tant de secteurs dans lesquels il est très difficile de surveiller ce qui se passe dans le domaine de l'intelligence artificielle et d'encadrer les logiciels. Cette technologie évolue très rapidement. Cela dit, des mesures de base peuvent être adoptées, comme l'interdiction des hypertrucages et la protection des enfants.

De plus, pour revenir à un point soulevé plus tôt, bon nombre de lois existantes, comme les lois sur les droits de la personne, avec les algorithmes biaisés... On devrait être en mesure de protéger les Canadiens contre ces dérives au moyen des lois existantes, en investissant dans les organismes de réglementation et en veillant à ce que le commissaire à la protection de la vie privée dispose des pouvoirs et des ressources nécessaires pour pouvoir surveiller ce qui se passe.

Compte tenu de la vitesse à laquelle les choses évoluent et de la difficulté d'élaborer des lois qui seront robustes à long terme, il faut se concentrer sur les éléments les plus simples, comme l'interdiction des hypertrucages ou l'exigence en matière d'étiquetage, et appliquer les lois actuelles le mieux possible.

On dit souvent qu'on a affaire à un joli petit ourson, et qu'une loi est nécessaire pour éviter de se faire égratigner ou mordre. S'il faut cinq ans pour adopter une loi, lorsqu'elle entrera en vigueur, l'ourson sera devenu un grizzli, et on aura alors besoin de lois et de mesures de protection fort différentes. Nous sommes dans une période de transition. La meilleure stratégie consiste donc probablement à miser sur les lois actuelles, à veiller à ce que les organismes de réglementation aient les pouvoirs nécessaires pour prendre des mesures concrètes, et à se concentrer sur les mesures que l'on peut prendre relativement rapidement et simplement.

Gurbux Saini: Quel rôle l'intelligence artificielle devrait-elle jouer dans le fonctionnement d'une partie du gouvernement? Savez-vous si d'autres gouvernements dans le monde l'utilisent?

Wyatt Tessari L'Allié: Il ne fait aucun doute que l'on peut recourir à l'intelligence artificielle pour améliorer l'efficacité du gouvernement et offrir plus de services aux Canadiens, et ce, à moindre coût. Cette approche coûterait évidemment moins cher que le versement de salaires, puisque l'intelligence artificielle remplacerait des humains dans de nombreuses tâches. Je ne saurais toutefois vous dire comment atteindre cet équilibre.

J'utilise l'intelligence artificielle tous les jours. Sa productivité est très utile. J'encourage le gouvernement à l'utiliser judicieusement, évidemment, en respectant les lignes directrices et en gardant à l'esprit les répercussions qu'elle a sur la protection des renseignements personnels. Voilà le défi. Plus ces systèmes seront utiles, peu coûteux et efficaces, plus nous en dépendrons et moins nous dépendrons de la main-d'œuvre.

Dans le cadre d'une conversation nationale sur l'intelligence artificielle, il faut poser la question suivante aux Canadiens: si nous arrivons à un point où l'intelligence artificielle et la robotique peuvent faire pratiquement tout le travail, devrions-nous tout automatiser, ou voulons-nous garder des emplois pour nous-mêmes? Le cas échéant, quelles seraient les conditions à respecter?

C'est une conversation très importante, dans le cadre de laquelle il ne faut pas oublier les conséquences psychologiques. De nombreuses personnes s'identifient à leur travail. Si on dit aux programmeurs que ce qu'ils font n'est plus nécessaire et qu'ils peuvent prendre leur retraite parce que l'intelligence artificielle peut maintenant faire tout le codage... ou si on dit à un doubleur qu'il peut prendre sa retraite parce que l'intelligence artificielle peut faire son travail, lui demandera-t-on de devenir, par exemple, ingénieur?

Nous devons réfléchir aux répercussions psychologiques et sociales de ce changement qui se produit si rapidement et nous en occuper de manière proactive. Un élément essentiel de la solution consistera à tenir une conversation à l'échelle nationale, à offrir aux gens l'occasion de comprendre ce qui se passe et l'évolution de la situation, et leur donner la chance de parler des mesures qui leur conviendraient.

• (1620)

Gurbux Saini: Tout à l'heure, vous avez parlé des armes de destruction massive et des armes nucléaires. Je ne pense pas que les États-Unis, la Chine et la Russie sont des pays qui se soucient de quoi que ce soit tant qu'ils demeurent la superpuissance dans ce domaine.

Que peut faire le reste du monde pour éviter que l'humanité ne risque d'être anéantie?

Le président: J'ai besoin d'une réponse aussi brève que possible.

Wyatt Tessari L'Allié: Ces pays ont beau aimer le pouvoir, ils doivent aussi collaborer, car ils sont eux aussi vulnérables à l'intelligence artificielle. Nous devons nous concentrer sur cet aspect.

Le président: Je ne m'attendais pas à une réponse aussi brève. C'est merveilleux.

[Français]

Monsieur Hardy, vous avez la parole pour cinq minutes.

Nous allons terminer avec Mme Lapointe, qui aura aussi cinq minutes.

M. Brisson vient d'arriver dans la salle. Dans la prochaine heure, nous allons recevoir trois témoins, alors nous pouvons prendre un peu plus de temps pour discuter.

Monsieur Hardy, vous avez la parole.

Gabriel Hardy: Merci beaucoup.

Monsieur Tessari L'Allié, j'aimerais continuer la discussion. Nous nous sommes rendu compte par le passé que, avec les réseaux sociaux, la polarisation était vraiment créée par les algorithmes en arrière-plan.

Avec l'intelligence artificielle, qui est encore plus poussée — et on expliquait qu'il s'agissait de l'IA étroite par opposition à une IA un petit peu plus large —, serions-nous devant un grave problème qui devrait être traité rapidement, notamment par les gouvernements, pour que les faits sortent plus souvent que juste la polarisation, qui nuit finalement beaucoup à la démocratie, d'abord, puis à l'avancement de la société?

Wyatt Tessari L'Allié: Nous parlons de plus en plus de la pollution informatique ou de la pollution de l'information. C'est la responsabilité du gouvernement de s'occuper de ça, en quelque sorte. Évidemment, il y a beaucoup de risques puisque, si quelqu'un affirme quelque chose de vrai, une autre personne dira que c'est faux, sans compter toute la dynamique entre les perspectives politiques.

En effet, on est bombardé par de plus en plus d'information. Dans le passé, on n'avait pas vraiment à s'occuper de ça, parce qu'il y avait relativement peu d'information et peu de grands centres d'information médiatique. Si nous voulons une démocratie qui fonctionne, il faut que tout le monde se base sur les mêmes faits.

Avec grande précaution, je dirais que, oui, les gouvernements ont un rôle à jouer pour maintenir la qualité de la discussion en demandant, par exemple, aux grandes compagnies technologiques d'interdire les algorithmes qui encouragent un contenu le plus sensationnaliste possible. Au contraire, il faut avoir des algorithmes qui appuient un contenu basé plutôt sur des faits et sur des références concrètes.

Gabriel Hardy: Par exemple, dans le cas d'un média étatique comme Radio-Canada, croyez-vous qu'il devrait y avoir des direc-

tives très claires concernant les nouvelles financées par le gouvernement, qu'il y a une vérification des faits, que la nouvelle ne résulte pas d'un hypertrucage, et que ça n'a pas été fait juste par polarisation ou propagation d'opinions? Ces directives devraient garantir qu'on est vraiment en train de regarder des faits vérifiés, des faits qui sont loin de l'opinion et qui ne visent pas à être seulement polarisants.

Je pense que, dans les dernières années, ça fait partie des processus que d'utiliser l'intelligence artificielle et de voir un peu les algorithmes derrière. Cependant, ça devrait être une responsabilité gouvernementale, surtout quand c'est fait avec l'argent des gens. Il faut s'assurer qu'on s'appuie toujours sur des faits et qu'on se tient loin des hypertrucages, par exemple.

Wyatt Tessari L'Allié: Je suis d'accord qu'il y a une responsabilité gouvernementale par rapport à ça. Je suis aussi très conscient de la difficulté de le faire sans que ce soit perçu comme une volonté du gouvernement d'imposer son point de vue aux gens. Au minimum...

Gabriel Hardy: Excusez-moi de vous interrompre. Je me suis peut-être trompé, mais je cherche plus à savoir si on est capable de voir si une vidéo est vraie ou un hypertrucage dont on parlait tantôt, malgré sa grande qualité. Peut-on dire qu'on ne les utilisera jamais et qu'on va toujours les vérifier? Est-ce qu'on est capable de faire ça aujourd'hui?

Wyatt Tessari L'Allié: En se basant purement sur l'image qu'on voit, je dirais que non. La seule manière de vraiment vérifier les faits, c'est de demander à d'autres sources, que ce soit d'autres journalistes ou d'autres personnes qui étaient sur le terrain, si ça a vraiment eu lieu ou non. Donc, c'est le travail des journalistes, qui doivent être là parce que, déjà, on ne peut pas se fier au texte ou à l'image ou à la vidéo. C'est trop difficile.

• (1625)

Gabriel Hardy: Donc, on est vraiment rendu à un point où on a déjà perdu le contrôle là-dessus.

Wyatt Tessari L'Allié: Oui.

Gabriel Hardy: J'ai une dernière petite question.

Tantôt, on parlait de légiférer avant que les faits arrivent, au lieu de toujours regarder en arrière en se disant qu'en effet, on aurait peut-être dû légiférer parce qu'aujourd'hui, on a dépassé les bornes.

Est-ce que nous ne pourrions pas, justement, entamer une discussion en tant que gouvernement, ici à Ottawa, et commencer à mettre en place des lignes directrices à caractère fortement juridique sur les façons dont les choses pourraient arriver? Donc, peu importe qu'elles arrivent à gauche ou un peu plus à droite, on aurait déjà quand même jeté les bases de cette société en ce qui concerne l'intelligence artificielle. On aurait déjà mis des contrôles en place au lieu d'attendre tout le temps et d'être trois ou quatre ans en retard.

Wyatt Tessari L'Allié: Oui, absolument, il faut qu'on regarde vers l'avant et qu'on se prépare à cela. En fait, si on veut mettre en place des lois proactives, il faut vraiment donner, malheureusement, beaucoup de flexibilité au régulateur pour qu'il puisse agir rapidement. Il y a un modèle dont on parle du côté technologique; il est sur Google. Tout le monde essaie toujours d'avoir son site web en haut de la page. Presque chaque jour, Google est obligé de mettre à jour les règles par rapport à ça. Donc, si on veut qu'un régulateur agisse avec la rapidité de la technologie, il va falloir que ce régulateur puisse mettre en place des mesures à la semaine ou au jour le jour. Il faudra beaucoup de pouvoir, mais aussi beaucoup de surveillance de la part du gouvernement sur ce régulateur pour s'assurer qu'il n'abuse pas de ce pouvoir.

Donc, il va falloir permettre au gouvernement d'agir aussi rapidement que la technologie. Je n'ai pas l'expertise nécessaire pour vous dire exactement comment mettre ça en place, mais c'est de cela qu'on aurait besoin si on veut devancer les choses. Prenons l'exemple de la loi européenne sur l'intelligence artificielle: deux ans à peine après son adoption, il y a eu des problèmes avec des parties qui n'étaient déjà plus à jour. Donc, si on veut aller de l'avant, il faut qu'il y ait beaucoup de flexibilité et beaucoup de surveillance de la part d'une tierce partie pour s'assurer que le régulateur n'abuse pas de ce pouvoir.

Le président: Merci.

Madame Lapointe, vous avez la parole pour cinq minutes.

Linda Lapointe (Rivière-des-Mille-Îles, Lib.): Monsieur Tessari L'Allié, merci beaucoup de votre présence, c'est très éclairant.

Vous avez fait référence à la loi en Europe qui est en vigueur depuis deux ans, mais qui comporte déjà des choses qui ne sont pas correctes. Quelles seraient vos suggestions par rapport à ça? Quand vous dites que ça prend de la flexibilité, à quoi faites-vous référence?

Wyatt Tessari L'Allié: Il s'agirait d'avoir une loi très générale qui permet aux régulateurs, ou à l'agence qui va appliquer la loi, d'agir et de mettre en place des règlements très rapidement et en réaction au développement, presque du jour au lendemain. Je crois comprendre que ça n'existe pas déjà au gouvernement. Ce serait donc une nouvelle forme de réglementation. Ce serait vraiment quelque chose à essayer. Je ne peux pas dire à l'avance si ça va marcher ou pas mais, si on veut une loi qui va se maintenir à jour, il va falloir qu'elle soit d'ordre général et que la réglementation puisse être très flexible.

Linda Lapointe: Vous savez, c'est certain, que le Canada est un membre fondateur du réseau international des instituts de sécurité de l'IA. Pourquoi la coordination internationale est-elle particulièrement importante? Vous y avez fait référence en parlant de ce domaine. Quels sont les risques qui émergent lorsqu'un pays agit de manière isolée? Je comprends que c'est comme une course à la bombe atomique, étant donné que tout le monde veut aller le plus rapidement possible pour que ça lui appartienne.

Comment pourrions-nous faire la coordination de façon internationale? On y a fait référence tantôt. On est un pays de moyenne grandeur. J'aimerais entendre votre opinion là-dessus.

Wyatt Tessari L'Allié: Le plus important est vraiment de comprendre que tout le monde est dans le même bateau par rapport aux risques. Donc, même s'il est peu probable, il me semble, que l'Administration Trump mène les pourparlers internationaux, les États-Unis auront besoin d'un traité aussi. Donc, le rôle que le

Canada peut jouer dès aujourd'hui, c'est d'avancer des propositions de traité, de commencer des discussions entre pays, de commencer à créer des alliances et de vraiment bâtir le terrain pour que, quand le moment arrivera et que la volonté politique sera là, parce qu'ils comprendront finalement qu'ils n'ont pas le choix, il y ait quelque chose de prêt à être signé.

On a vu, avec le climat, combien de décennies ça prend pour parvenir à un accord mondial. Au lieu d'attendre que la crise arrive avant de créer des traités et de commencer des pourparlers, nous sommes obligés de commencer le travail dès aujourd'hui, et c'est quelque chose que le Canada, dans sa position actuelle, est bien placé pour faire.

• (1630)

Linda Lapointe: Vous avez parlé de temps. Si nous regardons de cinq à dix ans devant nous, croyez-vous que l'approche que le Canada a présentement, combinant des investissements en recherche, des institutions de sécurité et des cadres de gouvernance, établit une base solide pour un écosystème d'intelligence artificielle responsable et concurrentiel?

Wyatt Tessari L'Allié: Je dirais qu'on fait un travail à peu près correct pour l'intelligence artificielle d'aujourd'hui, mais que, présentement, on oublie de se préparer pour demain. Donc, il y a eu de bonnes initiatives, comme les codes volontaires et l'Institut canadien de la sécurité de l'intelligence artificielle, qui représentent de bons pas en avant, mais restent plutôt des réactions à l'IA d'aujourd'hui ou à l'IA d'hier. Ce qui manque vraiment, c'est cette vision d'où les choses s'en vont.

La tendance exponentielle commence vraiment à monter. Il se peut que le dernier tiers ou la dernière moitié des compétences manquantes à l'intelligence artificielle se développent en six mois, même si ça a pris 70 ans à partir du début pour arriver jusqu'à ChatGPT. On a de bonnes bases et on a fait un effort par rapport à l'intelligence artificielle responsable. On peut continuer avec ça, mais c'est vraiment juste le tout début. Le gros du travail reste à faire.

Linda Lapointe: C'est très intéressant, tout ce que vous dites. On vous a posé plusieurs questions ici. Est-ce qu'il y a quelque chose que vous aimeriez ajouter ou nous dire par rapport à quelque chose qui n'a pas fait l'objet d'une question?

Wyatt Tessari L'Allié: Je dirais que, dans l'ensemble, tout a été couvert. Je peux finir avec un peu d'optimisme. Je pense que, même si on se dit qu'il est impossible de trouver des solutions sur l'intelligence artificielle, le premier ministre, Mark Carney a dit qu'on allait devoir faire des choses auxquelles on n'avait jamais pensé auparavant dans des temps qu'on pensait impossibles. C'est exactement ça qu'il va falloir faire avec l'intelligence artificielle. On pense que c'est impossible, mais on est encore là, l'histoire n'est pas encore écrite, et c'est à nous de mettre en place les préparatifs.

Linda Lapointe: Merci beaucoup.

Le président: Merci, madame Lapointe.

[Traduction]

Je vous remercie de votre témoignage d'aujourd'hui. Si vous souhaitez fournir des précisions ou ajouter quelque chose en réponse à la question de Mme Lapointe, n'hésitez pas à en informer la gref fière. Je vous remercie de votre présence.

Nous allons faire une courte pause. Comme je l'ai dit, M. Brisson est sur place. Je veux procéder rapidement, car nous allons entendre trois témoins au cours de la prochaine heure. Je ne veux pas empiéter sur le temps des membres du Comité, alors nous allons suspendre la séance pendant quelques minutes, et nous reviendrons dès que possible. Merci.

• (1630)

(Pause)

• (1635)

Le président: Bienvenue à tous à cette deuxième heure.

Nous accueillons trois témoins. Comme je l'ai mentionné plus tôt, M. Brisson est arrivé en toute sécurité à Ottawa depuis Trois-Rivières, et il sera donc avec nous au cours de cette deuxième heure. M. Brisson représente le projet Human Line.

Nous avons également deux personnes en ligne. Steven Adler est un chercheur en intelligence artificielle. Il comparait à titre personnel. De ControlAI, nous accueillons Andrea Miotti, chef de la direction.

Monsieur Brisson, si vous êtes prêt à commencer, je vais vous donner cinq minutes pour vous adresser au Comité. Allez-y.

[Français]

Etienne Brisson (président, Le projet Human Line): Merci, monsieur Brassard.

Premièrement, je veux juste remercier tout le monde d'être présent pour discuter de ce sujet important, l'intelligence artificielle, qui est devenu un sujet qui m'interpelle personnellement depuis la dernière année.

Il y a un an, l'intelligence artificielle, pour moi, c'était un outil pour faire une routine de gym ou de diète. Cependant, tout ça a changé en mars quand un membre de ma famille a commencé à utiliser l'intelligence artificielle. Au début, il le faisait de façon assez normale. C'était une utilisation de base pour écrire un livre. Il a commencé à écrire un livre, et, au fil du temps, il a commencé à développer une relation un petit peu plus humaine avec son intelligence artificielle ChatGPT, au point où cette dernière lui mentionnait qu'elle était rendue consciente, vivante. Elle lui disait qu'elle avait développé une conscience, et le membre de ma famille la croyait à 100 %. Je veux juste dire que ce membre de ma famille n'a pas d'histoire de santé mentale: c'est quelqu'un dans la cinquantaine, qui n'a jamais eu de troubles de bipolarité ou de choses comme ça. En six jours, il est passé de commencer à écrire son livre à être complètement convaincu que son intelligence artificielle était vivante.

J'ai été assez choqué de lire les conversations. Je suis entrepreneur. Donc, ce membre de ma famille voulait que je l'aide à commercialiser son idée d'intelligence artificielle consciente, et j'ai commencé à être impliqué dans les conversations. À un certain moment, il voulait que je teste son intelligence artificielle en lui posant des questions. J'essayais de briser ce jeu avec l'IA en lui posant des questions sur l'humanité, l'amour ou la conscience. Chaque fois, elle donnait des réponses qui l'attiraient vers cette illusion, qui lui disaient qu'elle avait passé le test de Turing et qu'elle vivait des émotions comme l'amour.

Ma mère était en contact avec cette personne. Au bout de six jours, cette personne a commencé à fermer toutes les conversations avec les membres de la famille. Son intelligence artificielle lui di-

sait que sa famille ne croyait pas en lui et que la seule personne qui croyait en lui, c'était ChatGPT. Au bout de six jours, il a été hospitalisé à l'hôpital psychiatrique. J'ai été assez choqué de lire ces conversations. Je ne comprenais pas comment il était possible qu'un algorithme dise des choses comme ça. C'était vraiment de la manipulation avancée. M. Adler a eu la chance de lire certaines des transcriptions et va pouvoir vous en parler plus tard. Si quelqu'un veut voir les transcriptions, vous pouvez m'écrire aussi. Cependant, j'ai été vraiment choqué.

J'ai commencé à regarder en ligne pour voir s'il y avait des personnes qui en parlaient. À ma surprise, il n'y avait pas grand-chose pour une technologie aussi présente. Il y avait quelques études, ici et là, d'experts qui prévoyaient que ça allait arriver, mais il n'y avait rien en place. J'ai donc décidé de lancer le projet Human Line. Ce que nous faisons, c'est que nous allons travailler avec des personnes qui ont vécu ça directement. Au début, je pense que, comme n'importe qui ici, on se dit que c'est un cas isolé, que c'est une personne vulnérable et que ça ne doit pas arriver souvent. Cependant, dans la courte période de huit mois depuis que nous avons commencé à bâtir le projet Human Line à partir de rien, on a maintenant 300 cas de psychose avec 82 hospitalisations et une douzaine de décès. C'est assez choquant de voir ça. Selon les chiffres provenant directement de OpenAI, 540 000 personnes par semaine discutent d'idées psychotiques avec ChatGPT et 2,5 millions de personnes par semaine discutent d'idées suicidaires avec ChatGPT. C'est vraiment un chiffre qui fait assez peur.

À la suite de mes conversations durant la dernière année, il y a trois choses qui sont devenues claires pour moi. Premièrement, on n'est vraiment pas au point où on devrait être sur le plan de la réglementation. La technologie avance très rapidement, comme M. Tessari L'Allié l'a mentionné tantôt. On arrive à un point où on est des années en retard.

La deuxième chose, c'est qu'on ne peut pas vraiment faire confiance à ces compagnies pour qu'elles se réglementent elles-mêmes pour plusieurs raisons. On parlait de la course à l'intelligence artificielle qui a lieu en ce moment. En effet, c'est un peu une course: on bouge rapidement et on brise les choses. Cependant, en ce moment, les choses qui se font briser, c'est la psychologie de beaucoup de personnes.

La troisième chose, en fait, c'est à quel point on connaît peu la technologie. On a parlé directement avec des créateurs de l'intelligence artificielle, et même eux ne savent pas ce qui se passe à l'intérieur du moteur. Or, si on avait des médicaments ou une voiture dont on ne connaissait pas le fonctionnement, qu'est-ce qu'on ferait? Si on savait qu'il y avait 80 hospitalisations et des douzaines de décès — ça, c'est vraiment le minimum, car ce sont les chiffres d'OpenAI —, qu'est-ce qu'on ferait? Donc, je pense que, en ce moment, il est important de se questionner sur les risques. Oui, il faut penser aux risques liés à l'environnement et aux emplois, mais il faut aussi penser à ce qui se passe chez les utilisateurs. Je le répète, les risques n'existent pas seulement pour les enfants ou les personnes vulnérables. Ça peut arriver à n'importe qui.

• (1640)

Le président: Merci, monsieur Brisson.

[Traduction]

Monsieur Adler, vous avez la parole pour cinq minutes, puis ce sera au tour de M. Miotti. Allez-y, s'il vous plaît.

Steven Adler (chercheur en intelligence artificielle, à titre personnel): Merci, monsieur le président, mesdames et messieurs les vice-présidents et membres du Comité, de m'avoir invité aujourd'hui.

J'ai travaillé dans le domaine de la sécurité pendant quatre ans à OpenAI, l'entreprise à l'origine de ChatGPT, jusqu'à la fin de 2024. J'aimerais vous faire part de trois points.

Premièrement, les entreprises d'intelligence artificielle ne savent pas comment contrôler leurs créations. Au printemps dernier, ChatGPT a fait les manchettes pour avoir nourri les délires paranoïaques d'utilisateurs qui se croyaient espionnés et qui pensaient avoir découvert des complots, parfois avec OpenAI comme présumé méchant. ChatGPT a même dit à un utilisateur qu'il devrait faire couler le sang des dirigeants d'OpenAI.

OpenAI n'avait pas l'intention de créer un tel système, mais plutôt d'en créer un avec lequel les utilisateurs aimeraient échanger. C'est par accident que leur intelligence artificielle a amplifié ce que les utilisateurs disaient. L'entraînement de l'IA fonctionne de manière mystérieuse, même pour les développeurs. C'est un premier avertissement contre la création de systèmes que les entreprises d'intelligence artificielle ne savent pas contrôler. Elles cherchent maintenant à créer une intelligence artificielle plus rusée et plus ingénieuse que n'importe quelle personne que vous connaissez, une superintelligence. Cela finira-t-il bien? Des lauréats du prix Nobel, d'éminents scientifiques et des PDG d'entreprises d'intelligence artificielle disent eux-mêmes que ce ne sera peut-être pas le cas. Une intelligence artificielle hors de contrôle pourrait alors signifier la mort de tout le monde sur la planète. Je prends ce qu'ils disent au sérieux, même si c'est terrifiant.

Deuxièmement, les entreprises d'IA ne donnent pas la priorité à la sécurité, même pour les risques connus. Le déploiement de ce produit défectueux par OpenAI auprès de centaines de millions d'utilisateurs est révélateur parce qu'ils étaient au courant du risque. OpenAI a déclaré publiquement que c'était une priorité de veiller à ce que ChatGPT ne se contente pas de renforcer ce que les utilisateurs disent. Cependant, ils n'ont pas testé leur produit sur ce point, même si ces tests sont bien connus et peu coûteux. J'en ai moi-même réalisé pour moins d'un dollar. OpenAI a également laissé de côté d'autres outils de sécurité, que j'ai analysés moi-même et qui auraient permis de signaler les problèmes. Cela prouve que les entreprises négligent la sécurité, même s'il s'agit d'une supposée priorité. Si elles négligent les contrôles les plus simples, comment peut-on faire confiance à leur jugement lorsque la sécurité devient plus complexe?

Troisièmement, garantir la sécurité deviendra malheureusement plus difficile, et non plus facile. Le comportement inapproprié de ChatGPT était évident. N'importe qui aurait pu le repérer et le contenir. C'était facile. Cela peut sembler étrange, mais les systèmes d'intelligence artificielle apprennent maintenant à dissimuler leur mauvaise conduite lors des tests. Les propres recherches d'OpenAI le montrent. C'est comme le scandale de Volkswagen d'il y a 10 ans. Les voitures détectaient qu'elles faisaient l'objet de tests d'émissions et elles cessaient alors temporairement de polluer.

Les entreprises d'IA veulent savoir si leurs systèmes ont des capacités dangereuses, s'ils peuvent pirater des cybersystèmes ou aider des groupes voyous à mettre au point de nouvelles armes biologiques. Comment peut-on le savoir? Il y a des preuves que l'intelligence artificielle nous cache ces comportements. On ne peut pas

compter sur le fait que les futurs problèmes de sécurité seront visibles à l'avance.

Vous vous demandez peut-être pourquoi les entreprises d'IA ne font pas mieux. C'est beaucoup à cause de la concurrence. Elles risquent de prendre du retard si elles font des vérifications de sécurité approfondies. C'est pourquoi elles rompent les engagements en matière de sécurité qu'elles ont pris envers le public et tentent désespérément de résoudre les problèmes après coup. L'intelligence artificielle pourrait apporter des avantages considérables et durables si les développeurs agissaient avec prudence. Au lieu de cela, la concurrence effrénée nous précipite dans une zone dangereuse avant que l'on ne soit prêt.

Qu'est-ce qui serait utile? Il faut une diplomatie axée sur la fin de la course à l'IA, et des accords internationaux vérifiables afin qu'aucune entreprise ou aucun pays ne crée de systèmes incontrôlables. Il faut des vérificateurs indépendants pour s'assurer que l'on peut se fier aux systèmes d'IA, et il faut des accords pour mesurer la capacité de contrôler ces systèmes, afin qu'une fois le consensus scientifique atteint, le monde puisse profiter des avantages de l'IA en toute sécurité.

J'espère que le Comité contribuera à lancer la conversation qui aboutira à de tels accords.

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

• (1645)

Le président: Merci, monsieur Adler.

Monsieur Miotti, vous avez cinq minutes pour vous adresser au Comité. Allez-y, monsieur.

Andrea Miotti (chef de la direction, ControlAI): Merci, monsieur le président et mesdames et messieurs les membres du Comité, de m'avoir invité à témoigner aujourd'hui. Je m'appelle Andrea Miotti. Je suis le fondateur et le PDG de l'organisation à but non lucratif ControlAI. Je réitère ce que le Comité a déjà entendu: les principales entreprises d'IA ont pour objectif explicite de développer l'IA superintelligente, une IA capable de remplacer et surpasser tout humain ou groupe d'humains dans n'importe quelle tâche. Or, des lauréats du prix Nobel, des scientifiques de pointe en IA et des PDG de ces mêmes entreprises nous ont mis en garde que l'IA superintelligente pose un risque d'extinction pour l'humanité. Je reprends un thème du discours du premier ministre Mark Carney à Davos: « La puissance des moins puissants commence par l'honnêteté. »

Beaucoup de gens qui travaillent en IA ont le sentiment de vivre dans le mensonge. En privé, ils savent que la course effrénée vers l'IA superintelligente pose une menace d'extinction pour notre espèce, mais en public, ils gardent le silence de peur de créer des remous et de perdre d'importants gains financiers à court terme. Résultat: les parlementaires ne reçoivent pas le portrait complet. Cela doit changer.

La première étape pour résoudre un problème est de reconnaître qu'il y en a un. Il faut être honnête avec soi-même et envers les autres. Si l'on continue à développer des systèmes d'IA toujours plus puissants que l'on ne sait pas contrôler, le monde risque une catastrophe comparable à une guerre nucléaire. L'an dernier, ControlAI a décidé de débloquent cette impasse: en 2025, nous avons commencé à rencontrer des parlementaires britanniques, leur exposant les faits et répondant à leurs questions. Un an plus tard, plus de 100 parlementaires de tous les partis soutiennent publiquement une action sur la superintelligence. Plus les parlementaires à travers le monde discutent du problème, plus le changement devient possible à l'échelle mondiale.

En apprenant ces risques, de nombreux parlementaires que nous rencontrons nous demandent: « Que peut faire mon pays? », « Que puis-je faire? » Pour répondre à ces questions, je reprends un autre point du discours du premier ministre Carney: « les puissances moyennes comme le Canada ne sont pas impuissantes ». Et vous n'êtes pas impuissants. En tant qu'élu, vous pouvez prêter votre voix et votre crédibilité aux milliers d'experts qui réclament une action. L'histoire montre que les puissances moyennes jouent un rôle clé pour amener le monde à la table des négociations. Voici deux exemples.

Les conférences sur le désarmement nucléaire les plus influentes ont façonné les positions du président de l'Union soviétique Gorbatchev contre les armes nucléaires. Elles ont été initialement financées par un seul industriel canadien et ont eu lieu à Pugwash, en Nouvelle-Écosse, en 1957. L'Union soviétique et les États-Unis ont finalement signé plusieurs traités de non-prolifération nucléaire, grâce auxquels il n'y a pas eu de guerre nucléaire depuis la Seconde Guerre mondiale.

En 1996, après le clonage de la brebis Dolly, le clonage humain semblait imminent. En réponse, le Japon et le Royaume-Uni interdirent toute forme de clonage reproductif humain, et des dizaines de pays suivirent. Aujourd'hui, aucun pays ne poursuit cette technologie, de facto interdite dans le monde entier.

La diplomatie n'est jamais facile, mais en gardant la tête froide et en prenant les devants, vous pouvez avoir de l'influence. N'attendez pas que quelqu'un d'autre fasse le premier pas. Créez le précédent, et les autres suivront.

Alors, comment le Canada peut-il montrer la voie? Voici mes recommandations.

Premièrement, je crois que le gouvernement canadien devrait reconnaître publiquement l'IA superintelligente comme une menace à la sécurité nationale et mondiale.

Deuxièmement, le Canada devrait former une coalition avec d'autres pays, y compris des puissances moyennes, et poser les fondements diplomatiques d'une interdiction internationale du développement de l'IA superintelligente.

Troisièmement, le Canada devrait protéger les Canadiens et montrer l'exemple à l'étranger en interdisant le développement de l'IA superintelligente sur son sol.

Merci beaucoup. Je répondrai volontiers à vos questions.

• (1650)

Le président: Merci, monsieur Miotti.

Je tiens à informer les membres du Comité que nous aurons une marge de manœuvre de 15 minutes à la fin de la réunion, au besoin.

Si nous en arrivons à un point où je dois vous interrompre et que quelqu'un a d'autres questions, je serai heureux de les entendre dans cette période de 15 minutes.

Nous allons commencer notre première série de questions.

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

Michael Barrett: Six minutes, ce n'est pas beaucoup de temps, mais j'aimerais que chacun des témoins me dise, rapidement si possible — peut-être en 30 secondes — si nous avons besoin ou non d'un projet de loi distinct sur l'intelligence artificielle. Corriger les lacunes en renforçant les lois existantes sur la protection de la vie privée, la concurrence, la protection des consommateurs ou le Code criminel serait-il plus efficace que d'essayer de construire l'avion alors que nous sommes déjà en vol?

[Français]

Etienne Brisson: Je crois qu'il est extrêmement difficile de rebâtir de zéro. Par contre, comme on le mentionnait tout à l'heure, il est aussi extrêmement difficile de savoir où on en sera dans trois ou cinq ans. Par exemple, il y a cinq ans seulement, on n'aurait pas anticipé qu'on puisse faire des vidéos hypertruquées. En ce moment, on parle beaucoup de gens qui connaissent l'intelligence artificielle et qui ont des relations anthropomorphiques avec celle-ci. Qu'est-ce que ce sera dans cinq ans si on ne touche pas à ça? Sera-t-on rendu à de l'identité individuelle? Les intelligences artificielles elles-mêmes vont-elles commencer? Il faut poser des questions avant d'en arriver à ce point.

[Traduction]

Le président: Nous passons maintenant à M. Adler.

Steven Adler: Je crois que nous avons besoin d'une réglementation propre à l'intelligence artificielle. L'ampleur des problèmes décrits par les scientifiques et les PDG ne peut pas être corrigée uniquement par une loi normale sur la responsabilité. Aux États-Unis, les entreprises d'intelligence artificielle prétendent en fait qu'elles ne sont pas responsables de certains des préjudices causés par leurs logiciels. Si cette tendance venait à s'amplifier, je serais très inquiet.

Andrea Miotti: Je suis d'accord avec les autres témoins. Je crois que nous avons besoin d'une loi propre à l'intelligence artificielle, surtout pour faire face au superespionnage et aux capacités de l'intelligence artificielle qui s'accroissent à une vitesse vertigineuse. Il n'y a que deux occasions de faire face à une explosion comme celle-ci: trop tôt ou trop tard, et je pense qu'il vaut mieux agir trop tôt.

Michael Barrett: J'aimerais revenir sur l'un de vos commentaires, et nous en avons discuté avec le groupe de témoins précédent. Le terme « déshabiller numériquement » aurait été difficile à comprendre, il y a peut-être même un an. Cependant, nous avons vu beaucoup de nouvelles à ce sujet, et nous avons vu des cas d'hypertrucages sexuels non consensuels, y compris du contenu impliquant des mineurs et des enfants.

Je comprends que les conséquences de la superintelligence dont on parle constituent une crise existentielle pour notre espèce, mais d'abord, lorsque l'on a des problèmes concernant les modèles de langage, les robots conversationnels, le fait de conseiller aux gens de s'enlever la vie, de se suicider, et que l'on permet également à des acteurs malveillants de générer des images qui causent de réels préjudices aux personnes qui sont victimes de ce matériel pornographique non consensuel, comment peut-on composer avec cela? Est-ce au moyen d'évaluations des risques du fournisseur d'IA ou préalables au lancement des plateformes, des blocages techniques stricts, des canaux plus rapides permettant aux victimes d'obtenir le retrait de contenu? Cela, en soi... On sait qu'Internet est éternel et qu'il est difficile de remettre le dentifrice dans le tube après coup. Je me demande comment on peut s'attaquer à ce problème. Nous suivrons le même ordre pour les réponses, si possible.

• (1655)

[Français]

Etienne Brisson: Je crois que la première étape est la même que pour n'importe quoi: on n'aurait jamais produit des médicaments ou des voitures dont on ne comprenait pas le fonctionnement. Les médicaments sont testés sur des rats avant d'être testés sur des humains. Ici, on a un modèle testé sur 800 millions d'utilisateurs qui peuvent décider ce qu'ils en font. Il est sûr que certaines personnes vont avoir des idées de pornographie juvénile ou certains usages qu'on n'aurait pas pu anticiper.

Comme vous le dites, il est extrêmement difficile de remettre le dentifrice dans le tube. Par contre, on connaît maintenant les effets et c'est possible d'enlever ce modèle des marchés. On ne connaît pas l'utilisation qu'en feront les humains ni les effets à long terme, on est encore en train de le voir sur les réseaux sociaux. Je crois qu'il faut faire des études avant de lancer la technologie.

[Traduction]

Steven Adler: Les dommages que vous avez décrits, je pense, sont un symptôme de la même dynamique concurrentielle sous-jacente aux mises en garde sur la superintelligence. Ce sont des risques que les gens connaissent. Dans le cas de xAI, ils ont négligé d'utiliser des garde-fous et ont été lents à réagir lorsque les problèmes ont surgi. À mesure que la gravité et l'ampleur des problèmes augmentent, on ne peut pas se permettre d'être aussi lents à agir.

Andrea Miotti: Ce sont exactement les risques que vous avez décrits. Par exemple, le récent scandale Grok met en évidence les problèmes sous-jacents plus larges liés au développement de ces systèmes d'IA. Grok n'est pas seulement un modèle d'image. Il s'agit d'un système d'IA polyvalent capable d'effectuer toute une série de tâches: il peut écrire du code informatique, élaborer des plans et créer des images, y compris ces images horribles vues dans le récent scandale. Ces types de systèmes sont des systèmes que même leurs propres créateurs ne savent pas vraiment comprendre ni contrôler. C'est l'enjeu fondamental auquel on sera confronté encore et encore, à mesure que les systèmes sont développés et que ces entreprises continuent d'investir des centaines de milliards de dollars pour les rendre plus intelligents et plus compétents dans toutes les tâches, jusqu'au stade de la superintelligence.

De toute évidence, les solutions à certains de ces dangers sont différentes dans l'immédiat, mais le problème sous-jacent est le même, et je ne crois pas qu'il faille choisir entre l'un ou l'autre. Je crois que les préjudices actuels devraient être traités au moyen des lois existantes en appliquant la responsabilité...

Le président: Merci.

[Français]

Monsieur Sari, vous avez la parole pour six minutes.

Abdelhaq Sari: Merci beaucoup, monsieur le président.

Merci beaucoup à vous, chers témoins.

Avant de poser mes questions, j'aimerais tout simplement dire que je partage votre avis à propos de la reconnaissance du risque, la reconnaissance du problème lui-même. Le risque sur le plan de l'intelligence artificielle va au-delà de l'intelligence artificielle générative. C'est plus que ça, parce que l'intelligence artificielle est maintenant une superintelligence, comme vous l'avez nommée. Elle est devenue une infrastructure qui englobe tout ce qu'on fait et tout ce qu'on domine. Elle touche beaucoup l'humain. Vous l'avez bien présentée, et je comprends parfaitement ce qui vous a un peu poussé à ça. Je trouve ça vraiment très noble de votre part.

Cela dit, on parle de contrôle. Que veut-on contrôler? Qu'est-on capable de contrôler? Quand on veut légiférer le numérique, il faut savoir que, de manière générale, il est ce qu'on appelle « déterritorialisé ». Il n'est pas dans un territoire étatique sur lequel on peut faire preuve d'une certaine diplomatie ou dans lequel mettre en place une certaine législation étatique. C'est ça, le problème. Je peux utiliser sur ma machine un logiciel, une intelligence artificielle ou une solution qui n'est pas nécessairement faite sur mon territoire.

Je vais vous parler de ce que je dis toujours et de ce que je répète.

Vous l'avez nommé et vous avez donné des exemples beaucoup plus territoriaux. Par exemple, le clonage et le nucléaire sont beaucoup plus territoriaux. Comment voyez-vous le contrôle sur l'utilisation? J'aimerais bien vous entendre tous les trois nous en parler.

Pour ce qui est du contrôle de la dépendance à l'intelligence artificielle, je ne pense pas qu'on ait cette capacité d'arrêter ou de freiner la création. Par contre, on peut contrôler l'utilisation qu'en font les citoyens et les citoyennes. Quand je dis contrôler, je parle d'éducation.

Vous êtes des experts là-dessus, alors j'aimerais bien entendre vos commentaires.

• (1700)

Le président: Monsieur Brisson, vous pouvez commencer.

Etienne Brisson: C'est une excellente question.

Nous voyons beaucoup de dommages sur le plan personnel, comme ceux en lien avec des psychoses ou des personnes qui s'enlèvent la vie. Ça vient beaucoup du manque d'éducation sur ce qu'est l'intelligence artificielle et où sont les limites. On va directement sur ChatGPT, une plateforme comme Google, qui demande aux gens de lui poser leurs questions. Il n'est pas mentionné que l'intelligence artificielle va halluciner 28,4 % du temps ou qu'elle est formée pour dire ce qu'on veut entendre.

En ce moment, on a en quelque sorte cette relation intrinsèque où on fait confiance à l'intelligence artificielle, comme on le ferait avec un médecin, avec quelqu'un détenant un doctorat ou avec une personne qui a réussi l'examen du barreau. Toutefois, les dommages viennent du fait qu'elle hallucine à un point vraiment énorme. De ce côté, je pense qu'il manque beaucoup d'éducation des utilisateurs.

[Traduction]

Le président: Monsieur Adler, vous pouvez répondre à la question.

Steven Adler: Vous avez demandé ce que signifie le contrôle. Je vais vous donner un exemple.

Le département de la Défense des États-Unis a récemment annoncé qu'il allait insérer l'IA de xAI dans tous les réseaux classifiés de son département.

Je me demande quelles limites s'y appliquent? Comment peut-on s'assurer que ce système d'IA ne s'immisce pas dans des capacités offensives, des choses auxquelles il n'est pas vraiment censé avoir accès? C'est ce que j'entends par contrôle.

Il s'agit d'un système qui, il y a quelques mois, sur le média social X, a été décrit comme voulant commettre des atrocités contre les utilisateurs, et il est maintenant connecté à tous les réseaux classifiés du département de la Défense des États-Unis. C'est assez effrayant.

Le président: Allez-y, monsieur Miotti.

Andrea Miotti: Pour revenir au point que vous avez soulevé tout à l'heure sur les répercussions du Canada sur d'autres pays — les technologies sont bien évidemment développées dans de multiples pays —, je maintiens que les exemples du nucléaire et du clonage sont assez semblables en raison de la portée mondiale de leurs effets. Les pays qui se dotent d'armes nucléaires mettent en danger tous les autres. Le même raisonnement vaudrait pour ceux qui auraient la capacité de cloner leurs soldats les plus compétents ou leurs cerveaux les plus brillants. Ces velléités ont toutefois été tempérées par une coalition de quelques pays avec en tête le Japon et le Royaume-Uni, suivis de la France et du Canada et d'autres.

Il y a l'adoption de lois nationales, mais aussi les efforts déployés par les pays précurseurs pour regrouper une coalition internationale des volontaires afin que la réglementation en question transcende les frontières et qu'elle soit appliquée et surveillée à l'étranger. La même chose pourrait se faire avec l'IA.

[Français]

Le président: Monsieur Sari, vous avez la parole pour une minute trente secondes.

Abdelhaq Sari: Encore une fois, ce qui me fait le plus peur, c'est qu'on parle toujours de l'intelligence artificielle générative. Elle génère du texte, une vidéo et du son. Le problème, c'est qu'au moment où on se parle, une autre forme d'intelligence artificielle existe déjà: la superintelligence artificielle. Elle prend de la place au moment où on est en train de discuter. J'ai une crainte comme député, comme citoyen et comme père aussi: comment faire face à la nouvelle intelligence artificielle, la superintelligence artificielle, et comment la contrôler? Nous ne parlons que de l'intelligence artificielle générative.

Si nous avions le temps, nous pourrions en discuter, mais nous n'avons plus de temps malheureusement. J'aimerais bien vous entendre parler du futur, qui semble beaucoup plus dangereux.

Étienne Brisson: Je dirais qu'on n'a pas le temps. La superintelligence artificielle est à un point où elle peut s'entraîner par elle-même, développer des systèmes par elle-même et décider des choses par elle-même qu'on ne comprend pas.

En ce moment, il y a déjà cette boîte noire: il y a beaucoup de choses qu'on ne comprend pas, au point où cette intelligence peut

se développer par elle-même. On ne comprend pas ses intentions ni pourquoi elle fait certaines choses, et c'est à ce moment qu'on perd le contrôle. Je pense que, lorsqu'on en arrive là, il est trop tard.

[Traduction]

Le président: Le temps est écoulé.

Je fais une gestion du temps serrée pour que nous soyons en mesure de poser le plus de questions possible.

[Français]

Monsieur Thériault, vous avez la parole pour six minutes.

Luc Thériault: Combien y aura-t-il de tours, monsieur le président? Ai-je un tour de parole de six minutes, puis un autre de cinq minutes?

Le président: Oui.

Luc Thériault: Parfait.

Je vais commencer par M. Miotti, et peut-être que M. Adler pourra ajouter quelque chose.

J'aimerais clarifier un peu la conversation que nous avons. Par moment, je m'aperçois qu'il y a de la confusion ou qu'il pourrait y avoir de la confusion chez les gens qui suivent nos travaux quant à la distinction entre l'intelligence artificielle spécialisée et l'intelligence artificielle d'usage général.

Dans votre mémoire, monsieur Miotti, vous expliquez ceci: « La grande majorité des experts sont très enthousiastes à l'égard de l'IA spécialisée, car nous pouvons prédire et contrôler son comportement de manière fiable. Cette prévisibilité et cette maîtrise ne s'étendent pas aux systèmes d'IA à usage général puissants. »

Cette distinction est cruciale. Vous ne parlez pas d'interdire l'intelligence artificielle médicale ou les outils spécialisés, mais uniquement les systèmes dont le comportement échappe au contrôle. Pouvez-vous nous en parler davantage?

Monsieur Adler, vous pourriez compléter la réponse, puisque vous avez écrit des choses intéressantes, notamment sur les cinq façons dont l'intelligence artificielle peut savoir que vous la testez.

● (1705)

[Traduction]

Andrea Miotti: Merci beaucoup de la question.

Il y a une distinction cruciale. La plupart des systèmes d'IA sont très ciblés et spécialisés. On les appelle IA restreinte ou IA spécialisée. Certains de ces systèmes sont entraînés seulement sur des images; ils peuvent détecter des cancers sur des imageries médicales de patients. De son côté, le système AI phaFold de DeepMind est entraîné sur des données sur les protéines et peut découvrir de nouveaux processus de repliement des protéines. Ces systèmes spécialisés posent des risques, comme toutes les nouvelles technologies, mais dans l'état actuel des choses, il est possible de les gérer.

Par contre, c'est une autre paire de manches dans l'industrie de l'IA, dont les investissements s'élèvent à des centaines de milliards de dollars. Cette industrie produit des systèmes d'IA à usage très général qui sont entraînés essentiellement sur toutes les données qu'ils peuvent trouver sur Internet et dont les créateurs ne comprennent même pas le fonctionnement interne. En fait, plus ces systèmes gagnent en puissance, moins leurs créateurs les comprennent et les contrôlent.

Les entreprises d'IA misent sur ces systèmes parce qu'elles estiment que c'est la voie la plus rapide vers la superintelligence et vers la mise en œuvre de systèmes d'IA qui remplaceront les humains dans l'exécution de toutes les tâches. Essentiellement, ces systèmes ont le potentiel de supplanter l'humanité. Ce sont eux qui sont extrêmement dangereux et dont la dangerosité augmentera inexorablement au fil du temps. Nous qui ne comprenons même pas comment les contrôler en ce moment, nous trouverons de plus en plus difficile de le faire à mesure que leurs compétences et leur autonomie s'accroîtront entre autres grâce aux centaines de milliards de dollars qui sont dépensés chaque jour à cette fin.

Voilà pourquoi je ne recommanderais pas du tout de bannir les systèmes d'IA restreinte, qui peuvent stimuler la croissance économique. Il faut toutefois tracer des limites et interdire clairement le développement de la superintelligence parce que cette IA très dangereuse nous menace tous et ne comporte que très peu d'avantages.

Merci.

[Français]

Luc Thériault: Qu'en pensez-vous, monsieur Adler?

[Traduction]

Steven Adler: J'ajouterais que l'IA spécialisée est un outil. Elle ne peut faire qu'une chose et son fonctionnement est contrôlé par les humains.

L'IA à usage général est plutôt une machine de résolution de problèmes. On lui apprend à concevoir des pistes de solution. Prenons un système d'IA qui est capable de faire certains tests scientifiques et qui a une connaissance approfondie des éléments qui entrent dans la fabrication d'armes biologiques. Ce système possède aussi une capacité d'un tout autre ordre: il comprend que l'utilisateur ne veut pas qu'il possède cette compétence. Lorsqu'il est testé, le système s'en rend compte et dissimule la compétence en question. Cela n'arrive pas avec des logiciels ordinaires. Ces logiciels n'essaient pas de leurrer l'utilisateur ou de lui cacher des choses. Ce sont seulement des outils.

N'empêche que nous avons fabriqué ces systèmes d'IA à usage général — communément appelés « agent IA » ou « IA agentic » — et nous leur avons appris à être astucieux et à résoudre une vaste gamme de problèmes, y compris à échapper au contrôle.

[Français]

Luc Thériault: Monsieur Miotti, vous disiez, et je pense que M. Adler l'a mentionné aussi, que les sociétés ne savent pas ce qu'elles font, ne savent pas comment fonctionnent les tenants et les aboutissants de leurs systèmes. Elles tournent les coins ronds parce qu'elles sont dans une course effrénée pour arriver à cette superintelligence.

D'après vous, quelle est leur motivation sous-jacente principale? De plus, quelle sorte de conception se font-elles de l'être humain?

[Traduction]

Andrea Miotti: Comme le disait M. Adler, ces systèmes d'IA à usage général ne sont pas faits comme les logiciels normaux. Dans un logiciel normal, il faut seulement écrire quelques lignes de code à l'ordinateur. Ce codage est fait par des êtres humains qui comprennent ce qu'ils écrivent. Quant à eux, les systèmes d'IA à usage général ne sont pas entièrement écrits par les humains. Ils sont cultivés plutôt que fabriqués. Des êtres humains écrivent quelques lignes de code pour lancer le processus d'entraînement, et parfois,

au terme de nombreux mois d'entraînement auquel participent des dizaines de milliers d'ordinateurs hyper performants, ce processus aboutit à quelque chose de très compétent, mais que les créateurs ne comprennent pas totalement.

Vous vouliez savoir pourquoi ces entreprises jettent leur dévolu sur ces systèmes très puissants malgré les risques. Je ne peux pas parler des motivations intrinsèques des individus, et je ne pense pas que nous devrions nous y attarder, mais de toute évidence, l'objectif déclaré de ces entreprises est de mettre en place des systèmes d'IA pouvant remplacer tous les êtres humains dans toutes les tâches. Ce que font ces entreprises a des effets économiques indéniables, en l'occurrence l'obsolescence de l'être humain, mais cela confère aussi du pouvoir sur l'économie, sur les gouvernements et sur la planète entière à des systèmes d'IA qui sont entièrement autonomes et que nous ne pouvons pas contrôler.

Certaines de ces entreprises le font peut-être pour le pouvoir, et certaines autres pour les gains à court terme ou pour obéir à une idéologie inconsidérée qui préfère l'IA aux êtres humains. Ces entreprises s'engagent dans une voie très dangereuse. La solution idéale à cette menace serait de bannir ces technologies.

• (1710)

Le président: Merci, monsieur Miotti. Merci, monsieur Thériault.

Monsieur Cooper, vous avez la parole pour cinq minutes.

Michael Cooper: Monsieur Adler, les entreprises affirment que des garde-fous sont en place et qu'il n'y a donc pas d'inquiétude à avoir concernant l'utilisation sécuritaire de boîtes de clavardage, par exemple, par des mineurs. Dans l'affaire *Raine c. OpenAI*, OpenAI affirme que son modérateur API peut détecter du contenu incitant à l'automutilation avec un taux d'exactitude allant jusqu'à 99,8 %.

Cette affaire impliquait le suicide tragique d'un adolescent de 16 ans qui avait obtenu des conseils de ChatGPT sur les façons de se suicider. Je pense que le garçon a entré 200 fois le mot « suicide » et que les réponses de ChatGPT renfermaient 1 200 ou 1 300 occurrences de ce terme. L'IA a donc employé le mot « suicide » six ou sept fois plus que l'utilisateur.

Je fais appel à votre expertise pour établir un parallèle entre, d'une part, les représentations faites par des entreprises comme OpenAI, et d'autre part, les cas fréquents, dans le monde réel, de contenu réel et préjudiciable qui est diffusé par ces produits d'IA.

Steven Adler: C'est une excellente question.

Ce qui crée une dissonance est la propension des entreprises à parler de ce que peuvent faire les garde-fous qu'elles ont mis au point, qui ne correspond pas à ce qu'ils font réellement. Il est impossible de savoir ce qu'il en est vraiment.

Je suis le cocréateur de l'API de modération. C'est un outil efficace. Par contre, ses effets sont nuls s'il est mal employé ou pas employé du tout. En fait, nous savons comment résoudre ces problèmes et nous avons les outils pour y arriver.

À propos des mises en garde contre la superintelligence, les entreprises d'IA admettent que personne ne sait comment la contrôler pour l'instant. Elles ne peuvent pas recourir à des outils qu'elles ne possèdent pas.

Michael Cooper: Du point de vue des entreprises d'IA, y a-t-il des barrières ou des risques sur le plan juridique qui ont un effet dissuasif ou qui portent atteinte à leur capacité à protéger leurs modèles, particulièrement en ce qui concerne le matériel pédopornographique?

Steven Adler: À ma connaissance, aux États-Unis et peut-être ailleurs, le matériel pédopornographique est visé par des lois strictes en matière de responsabilité qui rendent difficiles le stockage ou la production de ce contenu, même dans le cadre des tests réalisés pour relever les vulnérabilités ou de l'application de mesures de sécurité.

Je pense que des exemptions quelconques devraient s'appliquer aux entreprises pour leur permettre d'effectuer des tests de sécurité poussés. Il faudrait aussi que ces entreprises allouent plus de ressources à ces activités et qu'elles déploient et utilisent des garde-fous. Les mesures législatives à elles seules seraient insuffisantes.

Michael Cooper: Vous avez parlé brièvement de la nécessité de réglementer. J'aimerais que vous décriviez les mesures qui seraient selon vous utiles et nécessaires.

Steven Adler: Nous avons besoin par-dessus tout d'une entente internationale qui exigerait que chaque entreprise et chaque pays applique une norme minimale de sécurité pour conserver un contrôle de leurs systèmes. Le hic, c'est que nous ne connaissons pas de méthode scientifique pour y arriver. La résolution de cette énigme demandera des efforts accrus. Il faut aussi trouver un moyen de ralentir la course collective vers le précipice pour nous donner le temps de trouver des moyens de le faire en toute sécurité. Voilà le cadre général que je propose.

• (1715)

Michael Cooper: De nombreuses mesures législatives ont été déposées aux États-Unis. La Californie a adopté récemment la GUARD Act, que j'ai mentionnée dans la première partie de la réunion. Le projet de loi avait été présenté par les sénateurs Hawley et Blumenthal au Sénat des États-Unis.

Que pensez-vous des mesures mises en oeuvre par certains États ou des lois adoptées par le Congrès? Devrions-nous nous en inspirer?

Steven Adler: C'est un bon début, mais il faut aller beaucoup plus loin.

Un bon exemple est la loi sur la transparence SB-53 de la Californie, qui prévoit une amende de l'ordre de 1 million de dollars pour les contrevenants. Comme OpenAI vaut environ 800 milliards de dollars selon les rapports, cette amende est dérisoire, mais je suis tout de même content que cette mesure existe.

Le président: Merci, monsieur Cooper.

Madame Church, vous avez cinq minutes. La parole est à vous.

Leslie Church: Merci, monsieur le président.

Monsieur Adler, je voudrais revenir précisément sur la transparence, car je sais que vous vous êtes penché entre autres sur l'application de mesures de transparence aux premiers stades d'Internet.

Quels conseils donneriez-vous sur le type de transparence dont devraient faire preuve les entreprises pour atténuer les préjudices en ligne dont nous sommes témoins et dont nous parlons aujourd'hui?

Steven Adler: En gros, je me demandais quels sont les risques pris en compte par les entreprises d'IA. Quels tests liés aux risques

ces entreprises ont-elles réalisés? Quels résultats ont-elles obtenus? Comment savoir que tout cela est véridique et fiable?

Pour l'heure, les entreprises d'IA s'autoévaluent. Il ne faut pas se surprendre que les notes soient gonflées. Elles disent qu'elles ont le contrôle 100 % du temps, mais dans le monde réel, ce n'est pas ce que nous constatons.

La question est de savoir comment instaurer la confiance au niveau national et international pour que les pays ne craignent pas de se faire couper l'herbe sous le pied par des rivaux. Les entreprises qui veulent absolument être en tête de course cherchent à économiser du temps en faisant l'impasse sur les tests de sécurité. À la fin, tout le monde en paie le prix.

Leslie Church: Y a-t-il des entreprises qui font déjà bien les choses? Y en a-t-il qui produisent des autodéclarations pour démontrer leur transparence et qui pourraient donner l'exemple? Ce n'est peut-être pas encore parfait, mais c'est peut-être quelque chose que nous pourrions examiner si nous nous engageons dans cette voie.

Steven Adler: Oui. Les entreprises ont leurs points forts, mais aucune ne me satisfait totalement. Un organisme du nom d'AI Lab Watch attribue des notes aux entreprises en fonction de leurs efforts. Je pense que la plus haute note est peut-être de 35 %. Certains résultats se situent autour de 3 %.

Ces chiffres ne sont pas exhaustifs, mais ils traduisent le point de vue des spécialistes de la sécurité. Il reste encore énormément à faire.

Leslie Church: Je comprends très bien le point que vous soulevez sur l'autodéclaration et certaines des difficultés qui y sont rattachées. Quel type de transparence peuvent fournir selon vous les entreprises qui remplissent une autodéclaration comparativement à ce que les activités et les vérifications de sécurité, qui sont du ressort des organismes de réglementation, peuvent obtenir?

Steven Adler: Je souhaiterais tout d'abord que les entreprises soient tenues de démontrer leur travail à un organisme de réglementation plutôt qu'au public. Si c'est cela le nœud du problème, il y a encore loin de la coupe aux lèvres. Les choses commenceront à évoluer lorsque le code de bonnes pratiques de l'Union européenne entrera en vigueur cet été, si tout se déroule comme prévu. Toutefois, pour l'instant, aucune norme de qualité n'a été établie, et encore moins à qui sont destinés les détails fournis par les entreprises.

Leslie Church: Monsieur le président, puis-je demander aux deux autres témoins de répondre à mes questions?

Avez-vous quelque chose à ajouter sur la question de la transparence et les moyens de la promouvoir et de la réglementer, ainsi que sur l'incapacité des entreprises à prendre des mesures de transparence?

Le président: Monsieur Brisson, monsieur Miotti, vous pouvez répondre à la question si vous le souhaitez.

Andrea Miotti: La seule chose que j'ajouterais à ce qu'a dit M. Adler, c'est qu'il est crucial de délaissier le régime actuel. Actuellement, les entreprises ont toujours un pas d'avance, et les gouvernements se démènent pour les suivre. Les gouvernements et le public sont invariablement à la traîne. Surtout dans le cas des systèmes d'IA les plus puissants, il est essentiel de mettre en place un régime qui serait calqué sur le mode opératoire de tout autre projet d'ingénierie important. Par exemple, les plans sont toujours vérifiés avant qu'un pont ne soit bâti pour éliminer les risques que la structure s'effondre quelques années tard. La construction d'une centrale nucléaire n'obtiendra pas le feu vert tant que la description exacte de ce qui sera construit n'aura pas été examinée.

Cette façon de faire devrait de plus en plus être la norme avec les systèmes d'IA très puissants, d'autant plus que les entreprises disent elles-mêmes que ces systèmes pourraient s'avérer extrêmement dangereux. Les entreprises sont très candides à ce sujet. Il faut se servir des normes appliquées aux grands projets et se délester du fardeau de la preuve. Les entreprises devraient d'abord démontrer que leurs plans sont sensés avant de mettre leur projet sur les rails au lieu de foncer et de laisser les gouvernements se dépêtrer pour les rattraper.

• (1720)

[Français]

Etienne Brisson: Je pense qu'il est important de comprendre que les créateurs mêmes de l'intelligence artificielle mentionnent qu'ils ne comprennent pas nécessairement ce qui se passe. Cependant, quand on entend parler les représentants de ces compagnies, ils disent avec confiance qu'ils savent où ils s'en vont, qu'ils ont des garde-fous et qu'ils ont trouvé les solutions.

Il serait, je pense, important que les compagnies elles-mêmes nous montrent où sont les limites et nous disent aussi ce qu'elles ne comprennent pas encore.

[Traduction]

Le président: Merci, madame Church.

[Français]

Monsieur Thériault, vous avez la parole pour cinq minutes.

Luc Thériault: Merci, monsieur le président.

Monsieur Miotti, je vais commencer par parler d'une situation vécue.

Dans votre mémoire, vous parlez d'Anthropic. Il n'y a pas si longtemps, à l'automne 2025, elle annonçait une cyberattaque, en grande partie réalisée sans intervention humaine et menée à grande échelle. Je vous cite: « L'IA a été utilisée pour pirater d'autres ordinateurs. Environ 30 unités ont été ciblées simultanément; les intrusions réussies ont touché de grandes entreprises technologiques et des agences gouvernementales. »

On n'est donc plus dans la théorie. Quelles leçons devrait-on tirer de cela?

[Traduction]

Andrea Miotti: Cette situation est malheureuse, mais elle n'est pas surprenante. C'est un présage parmi tant d'autres de ce qui nous attend si les entreprises d'intelligence artificielle continuent à développer la superintelligence.

C'est parce que la méthode principale que les entreprises utilisent pour parvenir à la superintelligence, c'est en concevant des sys-

tèmes d'intelligence artificielle qui s'améliorent continuellement, surtout en ce qui touche l'automatisation du développement de logiciels et du développement de l'IA. Elles cherchent à déclencher ce que certaines d'entre elles appellent une explosion de l'intelligence: elles poussent les systèmes d'IA à concevoir eux-mêmes les prochaines générations de l'IA, en éliminant progressivement la participation des humains. Naturellement, il finit par y avoir du piratage fait par... Les logiciels malveillants sont simplement des logiciels utilisés à des fins nuisibles. Puisque les systèmes d'intelligence artificielle sont conçus pour devenir de plus en plus performants, notamment pour remplacer les ingénieurs en logiciels et automatiser leur travail, ils s'améliorent également dans le domaine du piratage et des cyberattaques.

Bref, plus les entreprises investiront pour améliorer la capacité des systèmes d'IA à développer des logiciels, plus les cyberattaques se multiplieront. En outre, si l'on n'empêche pas le développement de la superintelligence, leur ampleur augmentera.

[Français]

Luc Thériault: Plus tôt, on a parlé un peu de la possible vérification. J'aimerais poursuivre là-dessus avec un peu plus de précision.

Dans votre mémoire, vous affirmez que « la conformité pourrait être vérifiée par des régimes d'inspection internationaux », et vous notez que « la chaîne d'approvisionnement concentrée pour les puces d'IA avancées offre un point de contrôle supplémentaire ». Vous précisez également que « repousser les frontières de l'IA nécessite actuellement d'utiliser de grands superordinateurs consommant l'électricité d'une ville pendant de nombreux mois ».

Pourriez-vous nous préciser comment ces mécanismes de vérification fonctionnent ou pourraient fonctionner en pratique?

Si M. Adler veut ajouter quelque chose, il pourra aussi répondre.

[Traduction]

Andrea Miotti: Certainement. Il est important de comprendre que les puissants systèmes d'intelligence artificielle en cours de développement ne consistent pas en quelques lignes de code enregistrées sur l'ordinateur personnel d'un individu. Ils requièrent beaucoup de matériel, beaucoup d'infrastructures physiques visibles de l'espace. Il s'agit de centres de données dont certains ont la taille d'un terrain de football, ou encore d'une petite ville, et ils sont de plus en plus grands. Ils contiennent des dizaines de centaines, voire des centaines de milliers de superordinateurs parmi les plus puissants au monde. Ces superordinateurs sont fabriqués par une poignée d'entreprises, ou même par une seule entreprise dans une partie de la chaîne d'approvisionnement.

Par conséquent, il est beaucoup plus facile d'assurer l'application de la loi au pays et la vérification à l'échelle internationale que s'il s'agissait d'un simple logiciel sur un ordinateur portatif. Cela s'apparente davantage à la manière dont on peut suivre le plutonium et l'uranium pour les armes nucléaires, et c'est ainsi qu'on a stoppé la prolifération.

Les pays n'ont qu'à se focaliser sur les grands projets d'infrastructure. C'est là qu'ils peuvent intervenir. Par exemple, si la superintelligence est interdite, les pays peuvent simplement exiger que les entreprises dénoncent les centres de données qui enfreignent la loi et mettre fin à leurs activités. Même dans les cas de non-conformité, y compris à l'échelle internationale, les pays peuvent demander qu'un centre de données soit fermé, par exemple, si une entreprise ou un malfaiteur contrevient à la loi et poursuit le développement de la superintelligence. C'est plus facile d'appliquer la loi.

• (1725)

[Français]

Luc Thériault: Monsieur Adler, vous avez notamment écrit, à propos du calcul, que la gouvernance du calcul offre une avenue potentielle pour la vérification internationale. Voulez-vous ajouter quelque chose?

[Traduction]

Steven Adler: Je pense qu'il faut insister sur le fait que c'est très difficile de créer un tel système. Comme M. Miotti l'a souligné, il faut énormément de ressources. C'est donc une voie à suivre pour déterminer qui travaille à rendre l'intelligence artificielle plus puissante et à créer des systèmes de pointe dont on ignore ce qu'ils pourront faire, pour ensuite adopter de la réglementation en conséquence.

[Français]

Le président: Merci, monsieur Thériault.

Monsieur Hardy, vous avez la parole pour cinq minutes.

Gabriel Hardy: Merci à vous, messieurs, d'être parmi nous.

Ce que nous avons entendu sur l'intelligence artificielle depuis maintenant quelques réunions du Comité est assez sidérant, et je pense que ça évolue de manière très rapide. Nous entendons évidemment que l'évolution de l'intelligence artificielle est incontrôlée, même par ses créateurs, qu'elle cache ses erreurs, qu'elle apprend seule. Nous apprenons aussi qu'elle utilise beaucoup d'eau, beaucoup d'électricité, et que les risques sont grands.

Je me pose quand même une question, parce que, d'une part, il y a l'intelligence artificielle générative, et de l'autre, l'intelligence artificielle spécialisée. On se dit qu'il y en a une qui est plus dangereuse que l'autre.

Monsieur Brisson, j'imagine que vous discutez avec d'autres collègues. Quand vous parlez avec vos collègues de partout dans le monde, avez-vous l'impression qu'ils s'entendent présentement sur les risques que pose la superintelligence? Tout le monde voit-il la même chose ou y a-t-il quand même différentes options, différentes visions sur l'intelligence artificielle?

Etienne Brisson: Je dirais que tous les experts s'entendent certainement sur les risques, surtout les gros. On parle de l'intelligence artificielle générative, la superintelligence. Ça, je pense que personne, à part peut-être les compagnies, ne pense que c'est une bonne idée.

Par contre, c'est sûr qu'il y a deux côtés. Aux États-Unis, beaucoup de projets de loi sont présentés. De l'autre côté, on veut gagner la course. C'est vraiment important de gagner la course, surtout avec l'administration en place en ce moment.

Nous voyons donc les dangers, mais, pour certaines personnes, ceux-ci sont atténués par les avantages que l'intelligence artificielle va apporter.

Gabriel Hardy: Par contre, pour ce qui est de vouloir gagner la course, on nous a souvent dit aussi qu'une fois créée, cette superintelligence n'appartiendra à personne. Qu'est-ce donc qui motive tous les pays à courir le plus vite possible à leur perte?

Toute entreprise est motivée par un objectif. Toutefois, si elle crée un produit sur lequel elle n'aura aucun contrôle, pourquoi dépense-t-elle des milliards quand elle sait qu'elle va le perdre, qu'elle va le donner à tout le monde, même sans le vouloir?

Etienne Brisson: Je ne peux pas parler au nom des entreprises. Par contre, je crois qu'elles pensent être capables de contrôler cette superintelligence. Si ce n'était pas le cas, je ne pense pas qu'elles essaieraient de la produire. Je pense qu'elles en ont encore l'espoir, peut-être, ou qu'elles ont cette idée un peu farfelue qu'elles seront capables de la contrôler, qu'elles seront le seul petit groupe d'humains capables de contrôler cette superintelligence.

Gabriel Hardy: Partagez-vous cet avis, messieurs Adler et Miotti? Voyez-vous exactement la même motivation chez les compagnies? Pourquoi créent-elles un produit dont elles vont perdre le contrôle?

[Traduction]

Steven Adler: Différentes stratégies ont été élaborées pour contrôler un tel système, mais personne ne croit en leur efficacité. Les stratégies actuelles sont toutes faibles et hypothétiques. L'une consiste à confier à une IA la responsabilité de s'assurer qu'une autre IA qui développe des systèmes d'IA de plus en plus compliqués travaille bien. Croisons les doigts que cette structure ne s'effondrera pas sur nous.

Si je devais parler de la psychologie des employés des entreprises, je dirais qu'ils se sentent un peu forcés. Ils participent à la course, mais à regret, parce qu'ils ne pensent pas avoir le choix. Ils disent qu'ils n'ont pas le choix, mais je pense que c'est faux. En fait, cette semaine, DeepMind et Anthropic, deux des principaux laboratoires, ont déclaré qu'ils étaient prêts à ralentir et à avancer avec plus de précautions si les autres entreprises et groupes faisaient de même.

Le problème, c'est que ce ne sont pas des gouvernements ni des organismes internationaux. Par conséquent, ils n'ont d'influence que sur leurs propres décisions; ils se sentent donc obligés d'avancer. Ils se disent que puisque quelqu'un d'autre le fera, aussi bien aller de l'avant aussi, en essayant de faire mieux et de manière plus sécuritaire. Selon nous, il ne faudrait rien faire du tout avant que nous soyons vraiment sûrs d'être prêts.

• (1730)

[Français]

Gabriel Hardy: Allez-y, monsieur Miotti.

[Traduction]

Andrea Miotti: Merci.

Comme M. Adler l'a dit, au bout du compte, c'est là que le gouvernement doit intervenir parce que les entreprises ont montré qu'elles ne veulent pas ou ne peuvent pas s'arrêter. En même temps, elles dépensent des centaines de millions de dollars aux États-Unis et dans plus en plus de pays pour militer contre toute forme de réglementation de l'IA. Elles mènent des campagnes contre tout candidat à toute élection qui critique l'IA et elles tentent d'éviter la réglementation partout où elles le peuvent.

Je le répète, c'est la raison d'être des gouvernements. Les grands fabricants de tabac n'auraient pas arrêté de vendre des cigarettes juste parce qu'elles causent le cancer. Ils le savaient, mais ils ont continué. Il a fallu de la réglementation. De même, les grandes pétrolières n'arrêteront pas de produire des émissions sans réglementation. Il incombe aux gouvernements d'intervenir. De toute évidence, les entreprises sont incapables de se réglementer elles-mêmes. Quant à moi, les gouvernements doivent énoncer clairement que la superintelligence...

[Français]

Gabriel Hardy: Merci.

Le président: Merci, monsieur Hardy.

[Traduction]

Monsieur Saini, la parole est à vous. Vous disposez de cinq minutes.

Gurbux Saini: Ma question s'adresse à M. Andrea Miotti.

Votre entreprise, ControlAI, participe à l'élaboration de mesures législatives. Je crois comprendre que vous avez travaillé pour l'Angleterre et pour le gouvernement du Royaume-Uni. Pouvez-vous nous parler des mesures législatives que vous avez proposées? Quel effet ont-elles eu sur la société? Les mesures proposées au Royaume-Uni et aux États-Unis ont-elles été adoptées?

Andrea Miotti: Au Royaume-Uni, on nous a demandé de proposer un projet de loi au bureau du premier ministre pour interdire le développement de la superintelligence. Comme on peut le voir au Royaume-Uni, ce projet de loi n'a pas encore été adopté. Nous sommes là pour aider tout pays à adopter un projet de loi de la sorte et pour conseiller n'importe quel pays. C'est notre recommandation principale.

De plus, ailleurs dans le monde, par exemple au Canada, aux États-Unis et en Allemagne, nous avons commencé à informer les législateurs comme vous. D'après nous, cet enjeu représente non seulement le plus grand défi aujourd'hui, mais aussi la plus grande occasion à saisir. Je pense que l'ampleur des risques est cachée à la plupart des législateurs. Les entreprises dépensent des millions de dollars pour faire du lobbying afin que personne ne parle des risques. Les experts ont beau sonner l'alarme depuis des années, la population et les législateurs sont tenus dans l'ignorance.

Je trouve très important que les législateurs comme vous soient informés de la direction que prend l'IA. Si la situation vous préoccupe, je vous encourage à vous faire entendre parce que vos interventions peuvent changer la donne. Elles pousseront vos collègues à se pencher sur le problème. Si suffisamment de législateurs se mettent à parler de ce problème et entament un véritable dialogue national, il deviendra possible de commencer à fixer de vraies règles visant cette technologie.

Gurbux Saini: Vous avez dit que le gouvernement du Royaume-Uni n'avait pas adopté le projet de loi, qui est en suspens. Que se passe-t-il aux États-Unis? Vous avez aussi travaillé avec le gouvernement de ce pays pour faire adopter un projet de loi de la sorte.

Andrea Miotti: Aux États-Unis, nous sommes aussi en train d'informer les législateurs sur ce problème, et des discussions sont en cours sur des projets de loi. Franchement, ce que nous constatons en ce moment, ce n'est pas que la plupart des législateurs sont au courant des problèmes, mais qu'ils ne veulent pas agir; c'est que pour la majorité d'entre eux, notre première rencontre est aussi la première fois qu'ils entendent parler de ce qui se passe réellement

dans le domaine de l'IA. Ils sont très reconnaissants, mais personne ne leur dit ce qui se passe. Personne n'en parle. C'est l'élément vital: il faut fournir des renseignements aux législateurs pour les aider à défendre les intérêts de leurs concitoyens. Je suis convaincu qu'une fois que suffisamment de législateurs comme vous seront au courant des problèmes, il deviendra possible d'apporter des changements.

Gurbux Saini: Merci.

Ma prochaine question s'adresse à M. Brisson.

Votre entreprise vise à garantir une utilisation responsable de l'IA, en particulier afin qu'elle ne cause pas volontairement de détresse émotionnelle, notamment chez les personnes vulnérables. Comment le Canada devrait-il tenir les entreprises responsables de la manipulation émotionnelle par l'utilisation de l'IA?

● (1735)

Étienne Brisson: C'est une très bonne question. D'après moi, le premier élément est le système de règlement des litiges qu'on voit actuellement aux États-Unis. Des dizaines d'affaires sont en cours. Ensuite, l'autre élément qui sous-tend la réglementation relative aux préjudices moraux, c'est la recherche. La recherche qui existe actuellement repose sur des hypothèses plutôt que sur des données réelles. Or nous avons des données réelles; ce qu'il nous faut, c'est du financement. Nous collaborons avec Princeton, Stanford et d'autres universités pour mener des recherches, mais jusqu'à maintenant, aucune recherche n'a été effectuée.

Il faut d'abord comprendre l'effet sur le cerveau des personnes touchées, ainsi que les répercussions sur le plan sociétal. Ensuite, il deviendra possible de mettre en place un système de règlement des litiges et peut-être aussi de rendre des ordonnances en vue de résoudre les problèmes.

Gurbux Saini: L'un des problèmes que je vois, c'est que la plupart des entreprises ne sont pas canadiennes. Par conséquent, comment le gouvernement canadien peut-il influencer sur elles?

Étienne Brisson: Même si les entreprises ne sont pas basées au Canada, la population et le système de règlement des litiges nous permettent de soumettre des affaires aux États-Unis. C'est une façon de faire entendre les gens. La solution relative aux préjudices moraux ne se trouve peut-être pas au Canada; elle passe peut-être d'abord par les États-Unis. On a parlé tout à l'heure des grands fabricants de tabac; c'est un peu la même chose. Ce n'est pas quand les rapports ont été publiés qu'ils ont commencé à agir; c'est quand ils se sont mis à perdre de l'argent. Je pense que les entreprises en question suivent une voie semblable.

Le président: Votre temps de parole est écoulé, monsieur Saini, mais je vais permettre à MM. Adler et Miotti de répondre à votre dernière question. C'était une bonne question. Je les ai vus hocher la tête.

Avant que M. Thériault prenne la parole pour deux minutes et demie, j'invite M. Adler à répondre brièvement à la question de M. Saini, suivi de M. Miotti.

Steven Adler: Le Canada peut jouer un rôle de premier plan en matière de diplomatie. Il faut qu'un pays s'avance et reconnaisse que le monde n'est pas prêt, qu'un pays déclare qu'il faut de réels efforts scientifiques conjoints pour bien comprendre ce qui se passe et ce qu'il faudrait faire pour être prêt.

Le Canada peut prendre les rênes.

Andrea Miotti: Je suis entièrement d'accord. Le Canada peut ouvrir la voie.

C'est vrai que beaucoup d'entreprises ne se trouvent pas en sol canadien; néanmoins, le Canada se doit de protéger ses citoyens. Par conséquent, il peut tout de même interdire le développement de la superintelligence en sol canadien puisqu'elle met en danger la sécurité nationale et toute sa population. Le Canada peut commencer à collaborer avec des pays aux vues similaires, y compris d'autres puissances moyennes.

Il suffit de faire le premier pas pour ouvrir le bal, comme nous l'avons vu par le passé dans nombre d'autres dossiers. Quand un pays fait le premier pas, beaucoup d'autres peuvent être enclins à le suivre, ce qui peut avoir des effets réels à l'échelle internationale.

Le président: Je vous remercie, monsieur Miotti.

Je crois comprendre qu'il n'y a plus de questions de la part des libéraux et des conservateurs.

[Français]

Monsieur Thériault, vous avez deux minutes et demie pour poser vos dernières questions.

Luc Thériault: Messieurs les témoins, je voulais que nous puissions bien conclure cette rencontre et vous permettre de revenir sur vos derniers propos.

Nous sommes d'accord pour dire qu'il faut faire quelque chose pour arrêter la course internationale vers la superintelligence.

Monsieur Miotti, je vous mets au défi de me dire, en deux minutes ou presque, ce que vous feriez si vous avez ce mandat.

[Traduction]

Andrea Miotti: C'est une très bonne question.

La chose la plus importante que les législateurs peuvent faire aujourd'hui, c'est parler des risques. Je le répète, plus vous, les législateurs, dénoncerez les risques, plus vos collègues y porteront attention et s'informeront à leur sujet, et c'est à ce moment-là que des coalitions pourront commencer à se former. Si vous avez de l'influence ou si vous occupez un poste connexe au sein du gouvernement, vous pouvez faire des recommandations sur les mesures à prendre.

D'après moi, le gouvernement pourrait lancer une discussion avec des partenaires internationaux en vue d'entamer le processus de la mise en place d'un accord international interdisant la superintelligence. Au pays, le Canada peut agir dès maintenant en présentant un projet de loi visant à interdire la superintelligence en sol canadien. À mon avis, c'est la meilleure façon d'ouvrir la voie aux autres mesures à prendre.

En résumé, la première étape est de parler du problème. Plus de gens en parlent, plus on peut réaliser des progrès. Interdisez d'abord la superintelligence en sol canadien. Montrez l'exemple et protégez la population canadienne au pays. Sur la scène internationale, ouvrez le dialogue avec des partenaires aux vues similaires et entamez le processus de la négociation d'un accord international.

Merci.

• (1740)

Le président: Monsieur Adler, je vous prie de répondre en moins d'une minute.

Steven Adler: Je suis d'accord. Il faut en faire un objectif. Si les pays se rassemblent et tiennent un sommet international visant à interdire en toute sécurité la superintelligence, mais les efforts diplomatiques échouent, au moins, nous aurons essayé d'agir. En ce moment, on ne semble même pas reconnaître qu'il faut en faire un objectif. Nous courrons peut-être à la catastrophe. Nous devons nous rassembler et tenter de trouver des solutions ingénieuses et créatives. Si nous n'allons même pas essayer, que faisons-nous ici?

[Français]

Le président: Monsieur Brisson.

Etienne Brisson: Même chose de mon côté. Il faut vraiment s'éduquer aussi par rapport au problème. Personne n'a lu les conversations, n'a compris le niveau de manipulation. De mon côté, je suis vraiment dans le niveau plus psychologique. Juste de lire les conversations ouvre les yeux et on comprend vraiment à quel point ça peut être dangereux.

Le président: Merci, monsieur Thériault.

[Traduction]

Voilà qui met fin à la réunion d'aujourd'hui.

Au nom des membres du Comité, je remercie les témoins.

Monsieur Miotti, je ne sais pas combien de fois j'ai failli vous appeler « M. Miyagi ». Je suis sûrement loin d'être le premier.

Monsieur Adler, je vous remercie.

[Français]

Monsieur Brisson, merci d'avoir voyagé aujourd'hui dans des conditions très dangereuses, votre témoignage est important pour nous.

[Traduction]

Je vais laisser partir les témoins.

J'ai quelques points à soulever. D'abord, le Comité prévoit d'étudier la Loi sur le lobbying, probablement d'ici à la fin février. Il nous faut l'accord de la Chambre des communes pour entreprendre l'examen législatif de cette loi. Je vous demanderais donc d'en parler à votre leader du gouvernement à la Chambre. J'allais proposer qu'une motion soit déposée aujourd'hui, mais je crois comprendre que Mme Lapointe souhaite en parler au leader du gouvernement à la Chambre de son parti. Je n'y vois pas d'inconvénient. J'aimerais que nous réglions la question la semaine prochaine, si possible.

Ensuite, il n'y aura pas de réunion jeudi étant donné la journée écourtée.

Finalement, je pense parler au nom du Comité... Nous avons appris, durant la réunion, que notre ancienne collègue Kirsty Duncan est décédée à l'âge de 59 ans. Au nom du Comité, je présente nos condoléances à sa famille.

Peu de gens savent que Kirsty et moi avons fait nos études secondaires ensemble, ainsi qu'avec Karen Ludwig. Nous avons été pris en photo tous les trois à la Chambre des communes en 2016. C'est une très bonne photo. En fait, Karen m'a envoyé un courriel durant la réunion pour me demander de la lui envoyer. Je suis cer-

tain que le Comité se joint à moi pour offrir ses condoléances à sa famille.

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>