



CHAMBRE DES COMMUNES
HOUSE OF COMMONS
CANADA

45^e LÉGISLATURE, 1^{re} SESSION

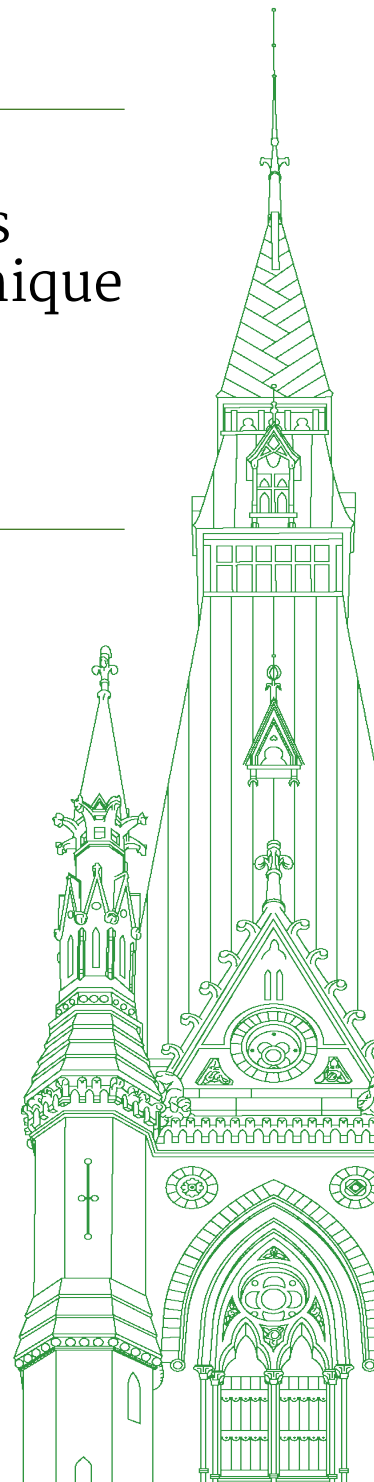
Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique

TÉMOIGNAGES

NUMÉRO 025

Le lundi 2 février 2026

Président : John Brassard



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

Le lundi 2 février 2026

• (1535)

[Traduction]

Le président (John Brassard (Barrie-Sud—Innisfil, PCC)):
Bonjour à tous. Je déclare la séance ouverte.

Je m'excuse auprès de nos invités pour notre retard. Nous avons eu des votes à la Chambre des communes.

[Français]

Je vous souhaite la bienvenue à la 25^e réunion du Comité permanent de l'accès à l'information, de la protection des renseignements personnels et de l'éthique de la Chambre des communes.

[Traduction]

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

J'aimerais souhaiter la bienvenue à nos témoins de la première heure de la séance d'aujourd'hui, qui sont tous présents par vidéoconférence.

De l'Institut pour l'avenir de la vie, nous avons Anthony Aguirre, directeur exécutif, et le professeur Max Tegmark. Nous recevons également David Krueger, qui est professeur adjoint au Département d'informatique et de recherche opérationnelle de l'Université de Montréal.

Nous allons commencer par M. Krueger. Viendront ensuite M. Aguirre, puis M. Tegmark. Chacun dispose de cinq minutes pour livrer son exposé au Comité.

Monsieur Krueger, vous avez la parole.

David Krueger (professeur adjoint, Département d'informatique et de recherche opérationnelle, Université de Montréal, à titre personnel): Bonjour. Merci de m'avoir invité. Je m'appelle David Krueger. J'enseigne l'apprentissage automatique à l'Université de Montréal et à l'Institut québécois d'intelligence artificielle, ou Mila.

En 2012, j'ai découvert l'apprentissage profond grâce aux cours en ligne de Geoff Hinton, et j'ai réalisé que cette nouvelle approche de l'intelligence artificielle, ou IA, pourrait aboutir à la création d'une IA superintelligente d'ici quelques décennies. Je suis allé étudier auprès de Yoshua Bengio, à Montréal.

À l'époque, je craignais déjà qu'une IA superintelligente puisse causer l'extinction de l'humanité, et je voulais savoir ce qu'en pensaient les experts. J'espérais trouver de bonnes raisons de ne pas s'inquiéter, mais ce que j'ai constaté, c'est que personne n'y réfléchissait vraiment. En fait, pendant la majeure partie de ma carrière dans ce domaine, le risque d'extinction de l'humanité par l'IA a été

considéré comme un sujet tabou, et les chercheurs craignaient pour leur carrière s'ils en parlaient.

Cela a malheureusement retardé de plusieurs années les débats publics cruciaux qui auraient dû se tenir sur la façon de gérer ce risque. Pourtant, depuis plus de dix ans, j'en parle chaque fois que j'en ai l'occasion. Je continue d'être consterné par les mauvais arguments avancés par certaines personnes pour éviter d'affronter le problème. D'un autre côté, j'ai pu constater avec le temps que de plus en plus de chercheurs s'inquiètent de cela.

En 2023, nous avons connu un tournant décisif et j'ai émis une déclaration où j'affirmais que la réduction du risque d'extinction causé par l'IA devait être une priorité mondiale au même titre que d'autres risques s'appliquant à l'échelle sociétale, comme les pandémies et la guerre nucléaire. Cette déclaration a été signée par de nombreux grands noms de l'IA, dont Hinton, Bengio et des centaines d'autres chercheurs actifs dans ce domaine. Malheureusement, en 2023, ChatGPT a également été lancé et les entreprises d'IA ont commencé à devenir extrêmement puissantes, ce qui est venu compliquer la réglementation de cette technologie.

La chose la plus importante que je souhaite souligner aujourd'hui, c'est que nous sommes toujours dans une situation où le monde ne prend pas encore les mesures nécessaires pour atténuer ce risque que l'IA conduise à l'extinction de l'humanité.

Quel est ce risque et d'où vient-il?

Les entreprises spécialisées dans l'IA tentent explicitement de créer des systèmes d'IA superintelligents. Il s'agit de systèmes qui seraient beaucoup plus intelligents que les humains dans tous les domaines et qui pourraient faire de manière autonome tout ce que les humains peuvent faire, y compris la robotique, et ce de façon bien meilleure, moins coûteuse, plus rapide, etc., que nous. L'objectif fondamental est de rendre les humains obsolètes et de leur prendre tous leurs emplois. Le hic, c'est que nous ne savons pas comment contrôler une IA superintelligente. En fait, nous ne comprenons pas comment fonctionnent les systèmes existants, car ils ont été développés, et non fabriqués, à l'aide de l'apprentissage profond. Malgré les milliers d'articles de recherche publiés au cours de la dernière décennie, cela reste un problème de recherche non résolu, et nous ne devons pas nous attendre à ce que des investissements, quels qu'ils soient, permettent de résoudre ce problème dans un avenir prévisible.

Nous ne savons pas non plus comment tester le caractère sécuritaire de ces systèmes d'IA. Les tests dont nous disposons peuvent montrer qu'un système d'IA est dangereux, mais ils ne peuvent pas montrer qu'il est sécuritaire. Là encore, nous ne pouvons pas nous attendre à ce qu'un investissement, quel qu'en soit le montant, permette de résoudre ce problème dans un avenir prévisible.

Au lieu de garder le contrôle de ces systèmes d'IA ou d'utiliser des pratiques rigoureuses pour en assurer l'« innocuité », les entreprises et les chercheurs en IA tentent de leur inculquer des objectifs et des valeurs particuliers afin qu'ils fassent ce que leurs concepteurs souhaitent. Sauf que la façon que nous utilisons pour faire cela est approximative, et même une petite erreur d'approximation pourrait conduire une IA superintelligente à réaffecter à ses propres objectifs les ressources dont nous avons besoin pour survivre. Aucune des approches dont nous entendons souvent parler, c'est-à-dire l'interprétabilité, les tests ou l'alignement, n'est techniquement adéquate. Bref, nous ne savons pas comment mettre au point une IA superintelligente en toute sécurité. À vrai dire, le plan consiste essentiellement à tenter le tout pour le tout.

Enfin, même si les entreprises qui développent une IA superintelligente ne perdent pas immédiatement ou rapidement le contrôle de cette dernière, nous devons tout de même nous attendre à un remplacement en règle des humains par l'IA, et ce, à l'échelle de la société. Cela signifie non seulement un chômage quasi total, mais aussi le transfert du pouvoir politique à des systèmes d'IA qui prendront des décisions trop rapidement pour que les humains puissent y participer de manière significative.

Je voudrais dire quelques mots sur les délais, car je pense que nous sommes dans une situation de crise aiguë. Si nous ne faisons rien, j'estime que nous sommes à environ cinq ans de la mise au point d'une IA superintelligente. Beaucoup sont d'accord avec cette projection. Nous devons immédiatement corriger le cap et travailler à empêcher la mise au point d'une IA superintelligente, et nous devons le faire à l'échelle internationale, ce qui prendra du temps. Nous ne pouvons pas nous permettre d'attendre d'autres preuves nous indiquant que le risque est imminent. Il existe déjà de nombreuses preuves que le niveau de risque lié à la voie sur laquelle nous sommes est inacceptable.

Dans mon pays natal, les États-Unis, le principal argument contre l'arrêt de la course à la fabrication d'une IA superintelligente est simplement qu'elle est inévitable, et que si nous ne le faisons pas, ce sera la Chine qui le fera. C'est faux. Un moyen simple d'y mettre fin serait de se débarrasser des puces informatiques d'IA avancées et des usines qui les produisent. Heureusement, la chaîne d'approvisionnement de ces puces est extrêmement concentrée, ce qui rend possible cette intervention ou d'autres interventions visant à contrôler et à limiter les moyens concourant à la production d'une IA superintelligente. Il existe peut-être des moyens moins coûteux, mais, compte tenu du risque, ce coût vaut la peine d'être payé si nécessaire. L'IA est une technologie qui n'a aucun précédent, et c'est l'avenir de l'humanité qui est en jeu.

• (1540)

Nous sommes en situation de crise. Nous devons agir immédiatement pour ralentir ou suspendre le développement de l'intelligence artificielle à l'échelle internationale. Cette question devrait être la priorité numéro un de tous les pays en matière de politique étrangère, y compris celle du Canada.

Je vous remercie.

Le président: Merci, monsieur Krueger.

Monsieur Aguirre, vous avez cinq minutes pour vous adresser au Comité. Veuillez commencer.

Anthony Aguirre (directeur exécutif, Institut pour l'avenir de la vie): Monsieur le président, membres du Comité, je m'appelle Anthony Aguirre. Je suis professeur de physique et président-direct-

teur général de l'Institut pour l'avenir de la vie, un organisme à but non lucratif international qui cherche à orienter les technologies transformatrices, en particulier l'IA avancée, dans des directions qui sauront profiter à l'humain.

Il y a 10 ans, lorsque nous avons fondé l'Institut pour l'avenir de la vie le domaine de la politique en matière d'IA n'existait pas et l'idée que l'IA puisse passer pour un humain de manière convaincante, prouver de nouveaux théorèmes mathématiques, écrire des codes informatiques complexes ou réussir presque tous les examens soumis habituellement aux humains relevait de la pure fantaisie. On pensait que cela prendrait des décennies, voire des siècles.

Aujourd'hui, l'IA est capable de faire toutes ces choses et les progrès ne montrent aucun signe de ralentissement. Si cette tendance se poursuit à ce rythme, c'est-à-dire si les progrès observés depuis 2003 se poursuivent au cours des deux prochaines années, l'IA pourrait bien être capable, d'ici un à cinq ans, d'accomplir toutes les tâches cognitives dont un être humain est capable. Les entreprises spécialisées dans l'IA visent à créer cette intelligence artificielle générale dans le but précis de remplacer les travailleurs humains plutôt que pour renforcer ces derniers ou leur prêter main-forte. C'est l'objectif que poursuivent ces entreprises. C'est ce qui sous-tend le modèle économique et les investissements colossaux dont nous sommes témoins.

Il ne s'agit pas d'un déplacement de main-d'œuvre comme pour les technologies que nous avons connues jusqu'ici, mais d'un remplacement complet de plusieurs dizaines de pour cent de la main-d'œuvre. Je pense qu'aucun pays n'a de plan viable pour faire face à un bouleversement aussi rapide et sans précédent du système de travail et de revenus humains, mais là n'est pas le plus grand risque. L'IA pourrait bientôt être capable de mener seule des activités de recherche et développement dans le domaine de l'IA. Elle pourra s'améliorer rapidement, progresser bien au-delà des capacités humaines et devenir une superintelligence. Il est ici question d'une IA qui rivalise non seulement avec les meilleurs humains, mais aussi avec l'humanité dans son ensemble. C'est en fait l'objectif déclaré de plusieurs entreprises spécialisées en IA. Plusieurs d'entre elles ont explicitement déclaré qu'elles cherchaient à mettre au point une IA capable de s'améliorer elle-même.

Je ne saurais trop insister sur le danger que cela représente. Ses développeurs ne seraient pas en mesure de prédire, de comprendre ou de contrôler très longtemps ce que ferait un tel système. Il fonctionnerait à une vitesse surhumaine, à une échelle surhumaine et avec une sophistication surhumaine vers des objectifs qui nous sont fondamentalement inconnus. C'est la raison pour laquelle tous ces éminents chercheurs et dirigeants d'entreprises ont signé la déclaration de 2023 où l'on affirmait que l'IA représentait un risque d'extinction pour l'humanité. Ce n'est pas un risque négligeable. Beaucoup de ceux qui développent ces systèmes estiment que les risques sont de 10 à 20 %, comme à la roulette russe. Certains experts qui ne participent pas au développement de cette superintelligence estiment que les risques sont encore plus grands que cela.

Je comprends qu'il est difficile d'accepter cette idée, étant donné son ampleur et son caractère incroyablement irresponsable, mais c'est pourtant ce qui se passe.

La bonne nouvelle, cependant, c'est que ce n'est pas inévitable. La voie tracée par quelques géants américains de l'IA n'est pas la seule possible. Il est encore possible d'opter pour une meilleure voie, pour celle qui consiste à développer des outils d'IA puissants et fiables qui renforcent et complètent les capacités des humains plutôt que de les remplacer, des outils dont le contrôle reste entre les mains des humains. Autrement dit, nous pouvons mettre au point des systèmes conçus dès le départ pour être fiables et contrôlables, et qui ont une portée et une fonction particulières, comme la détection du cancer, la prédiction du repliement des protéines ou les mathématiques avancées. Les outils d'IA de ce type permettraient des avancées médicales et scientifiques, d'énormes gains de productivité et la création de produits permettant aux gens de faire des choses qu'ils n'auraient jamais pu faire auparavant. Il suffirait de renoncer à l'objectif mal choisi de reproduire, puis de dépasser toutes les capacités humaines en travaillant sur des produits de substitution autonomes d'intelligence artificielle.

C'est pour ces raisons que l'Institut pour l'avenir de la vie a publié une déclaration demandant l'interdiction du développement d'une superintelligence, interdiction qui ne sera levée qu'après, premièrement, l'atteinte d'un large consensus scientifique sur sa mise en œuvre sûre et contrôlée et, deuxièmement, la constitution d'un vaste soutien au sein du public. Cette déclaration a été signée par des centaines d'experts du monde universitaire, par des employés de toutes les grandes entreprises mondiales d'IA, par des personnalités politiques, des chefs religieux et bien d'autres encore.

Je pense que le Canada est bien placé pour contribuer à cette réorientation vers une voie plus responsable que la voie actuelle. Pour ce faire, je recommande notamment au Canada de prendre les mesures suivantes.

Premièrement, soutenir les initiatives visant à créer des outils d'IA favorables à l'humain, qui visent à renforcer les capacités des humains plutôt qu'à les remplacer.

Deuxièmement, statuer que tous les systèmes d'IA doivent faire l'objet d'un contrôle humain significatif. Cela veut dire qu'il faut renforcer et clarifier les lois sur la responsabilisation afin que la responsabilité des actions entreprises par les systèmes d'IA incombe à qui de droit, c'est-à-dire aux utilisateurs, pour les systèmes sur lesquels ils exercent un contrôle important et, conjointement, aux développeurs et à ceux qui déploient si le contrôle exercé est insuffisant ou inexistant. Le contrôle ne doit jamais incomber aux systèmes d'IA eux-mêmes.

Troisièmement, il faut interdire la mise au point d'une superintelligence jusqu'à ce qu'il puisse être démontré qu'elle peut être contrôlée et utilisée de façon sécuritaire et qu'elle bénéficie d'un large soutien public.

Je vous remercie.

• (1545)

Le président: Merci, monsieur Aguirre.

Monsieur Tegmark, c'est à vous maintenant. Vous disposez de cinq minutes pour vous adresser au Comité. Veuillez commencer.

Max Tegmark (professeur, Institut pour l'avenir de la vie): Merci à vous, monsieur le président, ainsi qu'aux membres du Comité. Bonjour.

Je m'appelle Max Tegmark, je suis professeur au Massachusetts Institute of Technology et je mène des recherches sur l'intelligence

artificielle depuis de nombreuses années. Je suis également président et fondateur de l'Institut pour l'avenir de la vie, que M. Anthony a mentionné et qui est le plus ancien et le plus important organisme de surveillance de l'IA au monde.

La semaine dernière, j'ai eu le grand honneur d'écouter votre premier ministre, M. Carney, à l'occasion de sa visite à Davos. Son discours m'a beaucoup inspiré. Je vais suivre son exemple ici, à la fois pour dire la vérité à ceux qui sont au pouvoir, comme l'ont fait les professeurs Aguirre et Krueger, et surtout pour parler très honnêtement du monde tel qu'il est.

La situation actuelle en matière d'IA est vraiment démentielle. Aux États-Unis et au Canada, l'IA est moins réglementée que les sandwiches. Cela signifie que si j'ouvrais un café à Montréal, Toronto ou Ottawa, et que l'inspecteur de la santé publique venait me dire: « Vous avez 52 rats dans votre cuisine, alors je vais vous interdire de vendre des sandwiches pour le moment », je pourrais simplement lui répondre: « Ne vous inquiétez pas. En fait, je ne vais pas vendre de sandwiches, je vais juste vendre une petite amie pilotée par l'IA conçue pour s'entretenir avec les enfants de 12 ans. Je sais qu'il y a eu des problèmes avec le suicide d'enfants et autres liés aux robots conversationnels, mais j'ai un meilleur pressentiment à propos du mien. » Ce à quoi le représentant du gouvernement serait obligé de répondre: « D'accord, très bien », car c'est légal. Je pourrais dire que je ne vais pas vendre de sandwiches, que je vais simplement commercialiser la superintelligence, dont le professeur Krueger et le professeur Aguirre vous ont parlé, cette superintelligence dont les entreprises sont plus près la réalisation que la façon de la contrôler. Encore une fois, l'inspecteur ne pourrait dire que ceci: « D'accord, allez-y. C'est tout à fait légal. » En fait, il y a un homme au Canada qui a organisé un réseau où il a retrouvé plus de 300 victimes de robots conversationnels qui ont provoqué des suicides chez des enfants, des psychoses, etc. Il est absolument absurde que la vente de ce type de produits sans test préalable soit légale. Que peut-on faire à ce sujet?

Le Canada dispose actuellement de deux superpuissances. Vous avez un premier ministre et des dirigeants exceptionnellement audacieux, et vous disposez également de certaines des meilleures recherches en matière d'IA au monde, dans des endroits comme Montréal, Waterloo, Toronto et de nombreuses autres universités. La première façon de lutter contre ces cas de figure qui défient tout entendement — un exemple desquels on a pu voir la semaine dernière avec Moltbook, où plus d'un million d'agents IA ont écrit de manière autonome des messages sur la façon dont ils pourraient se débarrasser des humains et communiquer entre eux sans que ces derniers puissent comprendre ce qu'ils font — consiste simplement à commencer à traiter les entreprises d'IA comme serait traitée n'importe quelle autre entreprise dans n'importe quel autre secteur au Canada, soit en établissant des normes de sécurité. Avant de pouvoir lancer au Canada un nouveau robot conversationnel ou un compagnon animé par l'IA à l'intention des enfants, il faudrait soumettre ces produits à des essais cliniques rapides afin de s'assurer que les avantages l'emportent réellement sur les inconvénients, de la même manière qu'une entreprise pharmaceutique doit s'assurer qu'un nouveau médicament destiné aux enfants n'augmente pas les idées suicidaires.

Si vous faites cela, les entreprises américaines vont commencer à protester. La première chose qui se passera, c'est qu'elles protesteront et diront: « Alors nous devons quitter le Canada de façon définitive », tout comme OpenAI a menacé de quitter l'Union européenne si la loi européenne sur l'IA était adoptée. Sauf que vous devez adopter votre loi quand même, et ces entreprises resteront malgré cela au Canada, tout comme cela s'est produit avec OpenAI, qui est restée dans l'Union européenne. Cela vous protégera également contre tous les risques de perte de contrôle dont le professeur Krueger et le professeur Aguirre vous ont parlé, car votre loi interdira également aux gens de commercialiser des produits s'ils ne peuvent pas prouver qu'ils ne serviront pas à fabriquer des armes biologiques pour des terroristes ou à renverser le gouvernement canadien. Comme je l'ai dit, à l'heure actuelle, personne n'a la moindre idée de la façon de donner ce type de garanties sur les produits d'IA qu'ils fabriquent.

• (1550)

Au lieu de cela, le Canada connaîtra un âge d'or de l'IA, où des entreprises mettront sur le marché des outils pouvant être contrôlés de sorte qu'ils ne puissent plus pousser des enfants au suicide ou à d'autres comportements néfastes, comme c'est le cas avec l'industrie pharmaceutique canadienne, qui est très saine et très productive grâce à ses normes de sécurité.

Le président: Monsieur Tegmark...

Max Tegmark: Il y a donc moyen de régler cette problématique.

Merci beaucoup.

Le président: Merci, monsieur Tegmark. Nous avons légèrement dépassé le temps imparti, alors nous devons passer directement aux questions.

Nous allons commencer par M. Barrett, du Parti conservateur.

Allez-y, monsieur Barrett. Vous avez six minutes.

Michael Barrett (Leeds—Grenville—Thousand Islands—Rideau Lakes, PCC): Monsieur Aguirre, merci d'être des nôtres. Mes questions porteront sur votre déclaration préliminaire.

J'aimerais revenir sur votre description de la situation et vos recommandations au regard de la situation actuelle au Canada, où l'industrie — et le gouvernement — cherche à accélérer la commercialisation de l'IA. Vous avez signalé que les systèmes non sécurisés présentent des risques extrêmes. Pourriez-vous expliquer brièvement en quoi consisteraient, selon vous, les balises de sécurité minimales et non négociables que le Canada doit mettre en place avant d'autoriser le déploiement de modèles d'IA avancés?

Anthony Aguirre: Tout d'abord, selon un des arguments récurrents de l'industrie, il existerait une sorte de conflit entre l'objectif de rendre les systèmes d'IA sûrs, fiables et dignes de confiance et celui de les déployer efficacement dans l'économie pour en tirer des gains de productivité. Pour ma part, je pense que cela revient au même. Le principal obstacle à une meilleure adoption de l'IA dans l'économie et à des gains de productivité, c'est tout simplement la méfiance des gens. Ils n'ont pas le sentiment de pouvoir se fier aux systèmes d'IA, car ils n'en comprennent pas le fonctionnement.

Cela découle de la façon dont ces systèmes ont été conçus. On crée un système censé tout faire et remplacer entièrement les humains, au lieu de concevoir des outils précis pour des usages précis. Selon moi, si nous modifions cette approche pour concevoir des outils d'IA à finalité précise — par exemple, pour la recherche

scientifique, les mathématiques ou même des choses comme la gestion des calendriers — et si ces outils étaient réellement mis à l'essai pour veiller à ce qu'ils accomplissent ces tâches de manière correcte, fiable et sûre, cela représenterait un immense gain économique par rapport à la trajectoire actuelle, qui vise davantage à remplacer massivement la main-d'œuvre humaine.

L'élément crucial, c'est l'existence d'un système d'évaluation externe par l'entremise duquel un organisme indépendant ou une autorité indépendante peut vérifier si un système d'AI est sûr et fiable avant sa mise en marché, comme M. Tegmark l'a expliqué.

• (1555)

Michael Barrett: Oui, en effet. Les produits que nous voyons aujourd'hui ont été guidés par la demande. Je serais porté à croire que la perspective d'un retour sur investissement justifie les coûts de recherche et développement. C'est ce qui stimule la commercialisation des produits actuellement sur le marché, et c'est ce qui alimente aussi les discussions sur la façon d'aller plus loin. Si nous devions envisager une pause ou un changement de cap, quel levier proposeriez-vous pour amorcer cela?

À l'heure actuelle, monsieur Aguirre, il semble que les entreprises cherchent à offrir des produits et, par leur entremise, des solutions aux consommateurs — c'est-à-dire aux particuliers et aux entreprises — qui, franchement, y trouveront leur compte.

Le président: Monsieur Barrett, attendez un instant. J'ai arrêté le chronomètre.

Monsieur Tegmark, je vois que vous avez levé la main. Je suppose que c'est parce que vous avez quelque chose à ajouter. En règle générale, les questions sont adressées à la personne choisie par le membre du Comité. Lorsque ce sera votre tour de prendre la parole, vous pourrez peut-être parler de l'argument que vous voulez faire valoir, ou quelqu'un pourra vous interroger directement à ce sujet. La question de M. Barrett s'adresse à M. Aguirre.

Monsieur Barrett, je remets le chronomètre en marche.

Anthony Aguirre: Des produits d'IA adaptés aux besoins concrets des consommateurs — voilà exactement ce que j'aimerais voir, et je pense que c'est exactement ce que nous obtiendrons si nous nous concentrons sur des outils réellement conçus pour les gens afin de les aider.

Ce que nous observons, c'est un mélange entre ce genre de produits, dans certains cas, et une autre dynamique — la même qui a propulsé le déploiement des médias sociaux —, c'est-à-dire l'objectif d'essayer de s'emparer de la plus grande part de marché possible, le plus rapidement possible. Par exemple, je suis enseignant — nous le sommes tous les trois — et, du jour au lendemain, il est devenu impossible de donner aux étudiants des devoirs de rédaction ou même des problèmes de physique ou de mathématiques en raison de l'émergence soudaine d'un système d'IA apte à faire les devoirs à leur place. Les étudiants ont essentiellement l'impression qu'ils doivent utiliser l'IA pour rédiger leurs dissertations ou faire leurs exercices puisque les autres s'en servent. Ce n'est pas quelque chose qui a été réclamé par le système d'éducation. C'est quelque chose qui lui a été imposé, au détriment de nos étudiants.

Michael Barrett: Excusez-moi de vous interrompre, mais mon temps est limité. Je vous remercie d'avoir répondu à ma question.

Comment la solution ou la trajectoire que vous proposez devrait-elle se concrétiser? Est-ce une question de mise en application des lois? Est-ce une question de leadership au sein du gouvernement? Comment envisagez-vous cela? Comme vous l'avez dit, les entreprises cherchent à s'emparer de la plus grande part de marché possible, et ce, le plus rapidement possible.

Le président: Je vais vous demander de répondre en 30 secondes ou moins, s'il vous plaît.

Anthony Aguirre: Les compagnies pharmaceutiques veulent, elles aussi, maximiser leurs profits, mais elles sont quand même tenues de prouver que leurs produits sont sûrs et efficaces avant de les mettre en marché. Il devrait en être de même pour l'IA. Il faut d'abord préciser ce que le produit d'IA est censé accomplir, puis expliquer à une autorité externe comment le produit remplit la tâche de manière sûre et fiable, le faire approuver et, enfin, le mettre sur le marché. Je pense qu'une telle approche mènera à de bien meilleurs produits d'IA et à une bien meilleure adoption, ce qui changera la trajectoire actuelle: ainsi, au lieu de systèmes généralistes conçus pour remplacer les humains, on obtiendra des outils utiles conçus pour les aider.

Le président: Merci, monsieur Aguirre.

[Français]

Madame Lapointe, du Parti libéral, vous avez la parole pour six minutes.

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

Bonjour, soyez les bienvenus. Je suis heureuse que vous soyez avec nous.

Quelque chose me frappe, aujourd'hui. Plusieurs témoins, dont des chercheurs comme vous, sont venus nous parler de l'intelligence artificielle, et je dois dire que nous n'avons reçu qu'une seule femme. Quel est le point de vue des femmes sur le sujet, est-il différent de celui des chercheurs masculins qui sont venus nous dire qu'il y avait de grands problèmes? Vous avez parlé tantôt de la superintelligence et du fait qu'on ne sait pas comment la contrôler. Ça va affecter la vie de tout le monde: les femmes, les hommes et les enfants. Les chercheuses ont-elles une perception différente de celle que vous êtes venus nous exposer quant aux risques et aux dangers pour tout le monde?

Je vais demander à M. Krueger de répondre le premier. Ensuite, les autres témoins pourront répondre à tour de rôle.

• (1600)

[Traduction]

David Krueger: Je pense que cela reflète surtout la sous-représentation des femmes dans le domaine de l'IA en général. Un certain nombre de femmes partagent ces préoccupations, et vous pourriez également en discuter avec elles, notamment Tegan Maharaj, professeure à HEC Montréal, Ajeya Cotra, Katja Grace et d'autres. Il y a différentes perspectives dans le domaine en général.

Comme je l'ai mentionné, jecrois qu'avec le temps, le point de vue que je défends — à savoir qu'il s'agit d'un risque existentiel urgent — est devenu plus courant. Cela tient à une prise de conscience accrue de la problématique ainsi qu'à la rapidité des progrès observés dans le domaine de l'IA, qui a dépassé presque toutes les attentes.

[Français]

Linda Lapointe: D'accord. Merci.

J'ai une autre question à vous poser, monsieur Krueger.

Quel investissement public serait prioritaire pour réduire l'asymétrie entre les gouvernements et les grandes entreprises de l'intelligence artificielle?

[Traduction]

David Krueger: Si je comprends bien la question, il s'agit de savoir comment nous pourrions faire en sorte que le gouvernement possède des compétences techniques en matière d'IA comparables à celles des entreprises du domaine. À mon avis, c'est très difficile. La première étape serait plutôt d'essayer de responsabiliser d'avantage les entreprises et de les réglementer plus sérieusement.

Selon moi, si ces entreprises attirent autant de talents, c'est notamment parce qu'elles peuvent offrir des salaires beaucoup plus élevés grâce aux investissements considérables qu'elles reçoivent, dans l'espoir de créer une intelligence artificielle générale et une superintelligence pour ensuite, là encore, remplacer les emplois de tout le monde et en tirer des billions de dollars. À l'heure actuelle, le déséquilibre en matière de talents est immense et, en toute franchise, il ne se résorbera pas facilement. Il faudrait un processus assez long pour y remédier complètement.

Par ailleurs, l'un des principaux enjeux en ce moment est le manque de transparence dans les entreprises d'IA. Leurs activités ne font guère l'objet de surveillance. Il y a peut-être cinq ans, ces entreprises publiaient la plupart de leurs travaux de recherche. Le milieu de la recherche pouvait ainsi acquérir une compréhension commune du fonctionnement des systèmes les plus avancés. Cela permettait également une vérification plus efficace des systèmes de sécurité. Ce n'est plus le cas aujourd'hui. Les entreprises ne publient presque plus rien au sujet de leurs systèmes, et il n'y a aucune surveillance indépendante obligatoire.

Nous avons besoin de programmes permettant à des experts indépendants d'examiner les détails des systèmes d'IA. Il ne s'agit pas seulement d'évaluer ou de tester le modèle, mais aussi de disposer des renseignements nécessaires sur la conception du modèle et l'utilisation des données, lesquelles sont souvent protégées par des droits d'auteur. Il y a d'ailleurs eu un certain nombre de poursuites judiciaires d'envergure à ce sujet. Il faut également savoir quelles méthodes d'entraînement ont été utilisées. Le modèle a-t-il été entraîné d'une manière susceptible de provoquer délibérément une accoutumance, une dépendance et d'autres problèmes psychologiques? Nous savons que ces dérives se sont produites à maintes reprises dans l'industrie de la technologie, par l'entremise des médias sociaux et de l'IA. Enfin, comment le système est-il déployé, et que faut-il ajouter comme balises et mesures de protection supplémentaires? Comment l'entreprise surveille-t-elle les usages qui en sont faits? Les entreprises ont tout intérêt à comprendre comment leurs produits sont utilisés pour en tirer davantage de profits.

C'est quelque chose que le gouvernement doit aussi...

[Français]

Linda Lapointe: Merci, monsieur Krueger. Excusez-moi de vous interrompre. Je veux aussi poser une question à M. Aguirre, et je ne dispose que de six minutes de temps de parole.

Monsieur Aguirre, peut-on comparer les risques liés à l'intelligence artificielle avancée à ceux du nucléaire et du climat? On entend parfois parler de ça. Est-ce un phénomène fondamentalement différent?

[Traduction]

Le président: Veuillez répondre en 45 secondes, s'il vous plaît.

Anthony Aguirre: En ce qui concerne les similitudes et les différences, je pense que le risque particulier des armes nucléaires réside dans le fait qu'elles peuvent exploser. Il y a un parallèle à faire si nous créons une superintelligence et que nous en perdons le contrôle.

L'IA comporte déjà des risques importants et cause déjà beaucoup de torts parce qu'elle est très peu réglementée et qu'elle ne fait l'objet de pratiquement aucun test de sécurité. Comme M. Tegmark l'a mentionné, des preuves accablantes montrent que l'IA contribue de façon importante aux pensées suicidaires et à la dépendance aux médias sociaux chez les jeunes. Ces préjugés se manifestent déjà. Nous en voyons clairement les premiers effets sur la main-d'œuvre, avec le remplacement des humains par des systèmes d'IA. C'est ainsi que ces systèmes sont conçus; c'est leur objectif à grande échelle.

Je pense que le risque de l'IA, c'est qu'elle remplace les humains dans des rôles qui ne devraient pas être remplis par une machine — en tant que thérapeutes, compagnons, partenaires romantiques, travailleurs, etc. Cela peut aller jusqu'au remplacement des humains en tant que décideurs, et même au remplacement de l'humanité tout entière à plus long terme.

Je crois que les enjeux sont importants, mais ils peuvent prendre une forme plus diffuse.

• (1605)

Le président: Je suis désolé, monsieur Aguirre. Nous avons dépassé le temps imparti. Merci, monsieur.

[Français]

Je cède la parole à M. Thériault, du Bloc québécois, pour six minutes.

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

Je remercie beaucoup MM. Aguirre, Kruger et Tegmark de leurs témoignages éclairants et sans complaisance, qui donnent une raison d'être au volet éthique de ce comité.

C'est une réflexion fondamentale que nous lançons ici. Vous parlez de consensus scientifique à établir. Je pense qu'à chaque séance que nous tenons, ce consensus se dégage davantage.

Je vais commencer par M. Aguirre.

Dans votre article, qui demande de garder l'avenir humain, vous dites que la puissance de calcul peut être facilement quantifiée, comptabilisée et auditée avec relativement peu d'ambiguïté, une fois que de bonnes règles sont élaborées à cet égard.

Monsieur Kruger, vous décrivez l'intelligence artificielle de pointe comme un projet immense qui ne progresse que par un effort délibéré.

J'aimerais entendre vos deux points de vue sur la faisabilité technique d'un régime de vérification.

Monsieur Tegmark, vous pourriez ajouter des commentaires par la suite.

Je vous écoute, monsieur Aguirre.

[Traduction]

Anthony Aguirre: Je me ferai un plaisir de commencer.

Oui, je pense que c'est tout à fait faisable. Comme M. Krueger l'a laissé entendre, l'IA n'est possible que grâce à d'immenses quantités de calculs effectués au moyen de puces très spécialisées. Ces puces sont essentiellement fabriquées par une seule entreprise, à l'aide de machines construites par une seule entreprise et à partir de conceptions élaborées par une poignée d'autres entreprises.

Nous avons beaucoup entendu parler de cette capacité de calcul dans le contexte des restrictions et des contrôles à l'exportation, mais ces puces sont dotées d'un mécanisme de sécurité au niveau du matériel. Il ne s'agit pas seulement de savoir où les puces sont envoyées, mais aussi à quels contrôles elles peuvent être soumises. Ces puces intègrent des fonctions de sécurité matérielle qui permettent de vérifier leur niveau et leur type d'utilisation. Tout comme on peut désactiver à distance un téléphone volé, le matériel d'IA peut — et devrait — être configuré de telle sorte qu'il soit possible de le désactiver. À mon sens, plus l'IA devient puissante, plus elle nécessite un mécanisme de désactivation fiable.

Nous pouvons utiliser les capacités du matériel, c'est-à-dire la couche de base sur laquelle reposent ces systèmes d'IA, pour y ajouter des couches de contrôle et de vérification si nous définissons des lignes rouges à ne pas franchir à l'étape de leur conception et de leur déploiement.

David Krueger: Permettez-moi d'intervenir. Je suis d'accord avec M. Aguirre.

Il est très important de souligner l'ampleur des investissements dans ce domaine. Autrefois, les gens pensaient qu'il n'y avait pas vraiment de moyen de réglementer l'IA ni d'empêcher quelqu'un d'en faire un usage dangereux puisqu'il s'agissait d'un logiciel que n'importe qui pouvait utiliser sur son ordinateur portable, à même son gars. Ce n'est pas du tout le cas aujourd'hui.

Dans le contexte actuel, ces systèmes exigent des centaines de millions, voire des milliards de dollars. Ce n'est pas le genre de renseignements qui sont rendus publics en ce moment, mais les investissements sont énormes et continuent de croître, d'autant plus que le matériel est extrêmement spécialisé, comme M. Aguirre l'a dit. C'est là que se situe le principal levier d'intervention pour une réglementation internationale de l'IA, qui, comme je l'ai mentionné, s'avère absolument essentielle.

Je me contenterai d'ajouter qu'il est très important de réfléchir à la façon de rendre un tel régime aussi robuste que possible. La vérification pourrait, par exemple, prendre la forme d'une liste énumérant les types de systèmes d'IA autorisés à fonctionner sur les puces informatiques. Il pourrait également s'agir d'un mécanisme de géolocalisation afin que nous sachions où se trouvent les puces au cas où il faudrait les rappeler.

En fait, nous devrions cesser immédiatement la conception de systèmes d'IA plus puissants. La façon la plus sûre d'y parvenir serait d'arrêter de fabriquer les puces, plutôt que d'essayer de mettre en place un système de vérification technique plus compliqué et moins robuste. C'est mon point de vue personnel: pour nous donner un peu de répit, il faudrait cesser de fabriquer les puces et cesser de construire et d'entretenir les usines qui les produisent.

Comme les autres experts et moi-même l'avons mentionné, il nous reste peut-être quelques années. Ce n'est donc pas une situation où nous avons le temps d'essayer de trouver la solution parfaite. Nous devons mettre en œuvre immédiatement une solution qui ralentira ou interrompra l'incroyable rythme de progression vers la superintelligence.

• (1610)

[Français]

Luc Thériault: Êtes-vous d'accord là-dessus, monsieur Tegmark?

[Traduction]

Max Tegmark: Nous avons entendu des réponses techniques détaillées à cette question importante sur les mesures pratiques qui peuvent être prises. Je veux simplement ajouter qu'il ne s'agit pas ici — comme on l'a mentionné plus tôt — de suspendre les activités dans le domaine de l'IA. Je ne parle absolument pas de mettre l'IA sur pause. Je parle plutôt de mettre fin à certains usages: par exemple, des petites amies virtuelles pour des jeunes de 12 ans ou des systèmes d'IA pouvant apprendre aux terroristes à fabriquer des armes biologiques et d'autres produits qui sont clairement plus nuisibles que bénéfiques pour la population canadienne.

Ce n'est pas différent de ce que la Direction générale des produits de santé et des aliments de Santé Canada fait en permanence pour les médicaments. Ce n'est pas parce que Santé Canada interdit la mise en marché de produits pharmaceutiques n'ayant pas fait l'objet d'essais cliniques que le Canada doit suspendre la mise en marché des médicaments. Nous pouvons tout simplement appliquer à l'IA exactement la même approche que celle utilisée pour les médicaments, et ce serait un tout premier pas dans la bonne direction.

[Français]

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

Luc Thériault: Il y a quand même une différence à faire entre l'intelligence artificielle spécialisée et la superintelligence artificielle. Il y a des choses plus faciles à contrôler, comme les NIP de contrôle parental, que la course effrénée à la création de cette superintelligence.

Le président: Merci.

La parole est maintenant à M. Hardy, du Parti conservateur, pour cinq minutes.

Gabriel Hardy (Montmorency—Charlevoix, PCC): Je remercie tous les témoins d'être parmi nous.

Je vais commencer par M. Tegmark, que nous avons un peu moins entendu jusqu'à présent.

Monsieur Tegmark, vous avez souvent parlé du potentiel transformateur de l'intelligence artificielle, des objectifs qui peuvent être bénéfiques pour l'humanité.

Selon vous, le recours à l'intelligence artificielle dans le domaine médical fait-il partie des objectifs sur lesquels on devrait se pencher dans un avenir assez rapproché? À quel point l'intelligence artificielle pourrait-elle changer le domaine médical et l'améliorer?

J'oserais même vous demander si les gouvernements devraient se joindre à l'aventure afin d'aider notre société à guérir des maladies chroniques ou dégénératives.

J'aimerais que vous nous parliez du domaine médical et de l'intelligence artificielle.

[Traduction]

Max Tegmark: Je vous remercie. Il s'agit en fait d'un enjeu qui me tient particulièrement à cœur.

L'IA recèle un énorme potentiel pour améliorer les traitements médicaux, guérir le cancer, et ainsi de suite. Il ne s'agit pas d'une perspective lointaine, mais bien d'une réalité au moment où je vous parle. Même si toutes les entreprises spécialisées dans l'IA parlent beaucoup de guérir le cancer, une seule d'entre elles est parvenue à réaliser des progrès concrets. Il s'agit de Google DeepMind, qui a lancé AlphaFold, contribuant de manière concrète à la découverte de médicaments, ce qui lui a d'ailleurs valu le prix Nobel de médecine.

Plusieurs formes de cancer sont passées d'un taux de mortalité de 80 % à un taux de seulement 20 %; nous avons donc parcouru environ 80 % du chemin vers la guérison du cancer. Le principal risque qui m'inquiète est simplement que nous gaspillions tous ces incroyables avantages en laissant les médicaments basés sur l'IA totalement non réglementés, ce qui pourrait provoquer un mouvement de ressac. Beaucoup d'entre vous se souviennent d'un produit appelé thalidomide qui était commercialisé au Canada et aux États-Unis aux femmes enceintes souffrant de nausées matinales. Comme l'industrie pharmaceutique n'était alors absolument pas réglementée, cela a mené à la naissance de plus de 100 000 bébés sans bras ni jambes en Amérique du Nord, ce qui a conduit à la création de la Food and Drug Administration.

Oui, laissons les entreprises innover, de manière incroyable, et guérir des maladies grâce à l'IA, mais traitons-les de la même manière que nous traitons les entreprises pharmaceutiques et veillons à ce qu'elles ne commercialisent pas leurs produits avant qu'ils aient été correctement testés.

[Français]

Gabriel Hardy: Ce que vous dites, en gros, c'est que l'intelligence artificielle peut vraiment aider la planète à régler énormément de maladies chroniques, même dégénératives, mais que le problème n'est pas tant la vitesse de la progression, que le fait de laisser des compagnies uniquement motivées par le profit mettre sur le marché des produits qui ne sont pas encore adaptés.

Par contre, je crois que ça pourrait permettre une grande avancée dans les traitements ultraspecialisés. Voyez-vous aussi dans vos recherches ou dans l'ensemble du marché de l'intelligence artificielle qu'il est possible d'avoir des traitements encore plus pointus et adaptés à chaque être humain au lieu d'avoir des approches plus générales?

• (1615)

[Traduction]

Max Tegmark: Tout à fait. Il existe un énorme potentiel dans la mise au point de traitements personnalisés qui permettent, par exemple, de séquencer l'ADN d'un patient atteint d'un cancer, d'identifier les mutations spécifiques présentes dans ses cellules cancéreuses et de développer un traitement sur mesure, adapté à son cas particulier. Les possibilités dans ce domaine sont absolument incroyables.

Je crois fermement à l'innovation et à l'innovation dans le secteur privé, et la clé pour y parvenir est de mettre en place les incitations adéquates. Dans l'industrie pharmaceutique, l'aviation et même la restauration, les entreprises innoveront pour fabriquer des produits sûrs dont les avantages l'emportent sur les inconvénients, car ce sont les seuls qu'elles sont autorisées à vendre. Si nous parvenons à mettre rapidement en place les incitations adéquates pour le secteur de l'IA, alors ces entreprises qui, à mon avis, agissent actuellement de manière très imprudente, se concentreront rapidement sur l'innovation afin de se démarquer en proposant des produits sûrs. Je ne blâme pas les entreprises, mais bien le manque d'incitations adéquates à leur égard.

[Français]

Gabriel Hardy: Merci beaucoup.

Monsieur Krueger, on dit depuis tantôt — je crois que ce n'est pas la première fois qu'on l'entend — qu'à un certain moment, les entreprises privées vont même envisager de retirer les humains de l'équation, faire en sorte qu'il y en aura de moins en moins parmi leurs employés, pour les remplacer par l'intelligence artificielle.

Est-ce déjà en place? Les entreprises font-elles des études pour connaître le rendement qu'elles pourraient atteindre? Voyez-vous de manière générale les grosses entreprises de technologie se dire qu'elles ont bien fait de remplacer un humain par l'intelligence artificielle ou si, dans leurs calculs, elles se rendent compte que de prendre en charge, de vérifier et de contrôler ces fameux nouveaux robots d'intelligence artificielle est aussi coûteux que d'avoir un humain qui fait bien le travail?

[Traduction]

Le président: Monsieur Hardy, votre temps de parole est écoulé.

[Français]

Monsieur Hardy, veuillez m'excuser. Nous allons poursuivre avec Mme Church pendant cinq minutes, car je pense que la réponse à votre question sera trop longue.

[Traduction]

Madame Church, à vous la parole, je vous prie.

[Français]

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

C'est une bonne question, monsieur Hardy.

[Traduction]

Ma question s'adresse à M. Tegmark.

Vous avez beaucoup parlé des normes de sécurité contraignantes, et j'apprécie la comparaison avec notre approche en matière de réglementation des médicaments et des produits pharmaceutiques.

À titre de parlementaires et de législateurs, par où pensez-vous que nous devrions commencer? C'est une chose d'examiner certains des résultats de l'utilisation des robots conversationnels et de l'IA, en particulier lorsqu'ils touchent à la sécurité des enfants. Je pense que ce sont là des domaines très importants et évidents dans lesquels nous devons réfléchir à la manière d'assurer la sécurité de la population, et notamment des enfants.

Comment pouvons-nous, d'une certaine manière, saisir l'étendue du fonctionnement de l'IA dans tant de domaines différents, qui touche de nombreuses industries et de nombreux domaines poten-

tiellement problématiques? Comment pouvons-nous intégrer cela dans un modèle réglementaire qui nous permettrait de lutter efficacement contre certains des préjudices que vous venez de soulever?

Max Tegmark: Il s'agit d'une excellente question.

Pour simplifier les choses, on peut dire que les divers usages de l'IA suivent simplement la même approche que celle que nous avons dans toutes les autres industries puissantes, à savoir qu'il incombe à l'entreprise d'innover et de démontrer à des experts indépendants nommés par le gouvernement que les avantages l'emportent sur les inconvénients.

Je commencerais plutôt par un enjeu essentiel de sécurité publique, à savoir la sécurité des enfants. Aux États-Unis, environ 95 % des républicains et des démocrates s'accordent à dire que le gouvernement fédéral doit légiférer en ce sens. On parle d'une vaste coalition allant de Bernie Sanders à Steve Bannon, et je pense que nous verrons probablement un projet de loi en ce sens aboutir chez nos voisins du Sud.

Une fois établi le précédent selon lequel nous allons traiter l'IA comme n'importe quel autre secteur industriel, nous pouvons ajouter à la liste des normes de sécurité non seulement l'interdiction d'augmenter considérablement le risque de suicide chez les enfants, mais aussi les enjeux de sécurité nationale. Par exemple, vous ne pouvez pas vendre des produits qui pourraient permettre à des terroristes de fabriquer des armes biologiques. Vous ne pouvez pas commercialiser des produits qui pourraient être en mesure de renverser un gouvernement, comme nous l'ont expliqué les professeurs Aguirre et Krueger. Cela découle naturellement de cette approche simple qui consiste à traiter les entreprises d'IA comme n'importe quelle autre entreprise.

Je voudrais ajouter une chose. Si cette idée de perte de contrôle semble étrange, elle est en réalité très évidente et remonte à Alan Turing dès 1951: si l'on assemble un groupe de robots beaucoup plus intelligents que tous les humains, alors bien entendu, ces robots seront capables d'assembler des usines afin de fabriquer de nouveaux robots encore plus sophistiqués. C'est exactement ce que les entreprises essaient d'accomplir.

Par ailleurs, comme l'IA peut fabriquer davantage de robots, elle répond d'une certaine façon à la définition d'une espèce. Rendez-vous au zoo le plus proche et demandez-vous qui se trouve actuellement dans les cages. De quelle espèce s'agit-il? Observez-vous des humains? Non, ce n'est évidemment pas le cas. Pourquoi? Parce que nous sommes l'espèce la plus intelligente sur Terre, et je suis convaincu que cela doit demeurer ainsi. Ne laissons pas les entreprises vendre un...

● (1620)

Leslie Church: Permettez-moi d'intervenir ici. Je pense que vous percevez probablement notre vif intérêt pour la manière dont nous abordons cette question.

Permettez-moi de m'adresser un instant à M. Aguirre, car vous avez tous deux participé à l'Institut pour l'avenir de la vie.

Je suis très intéressé par le concept d'IA comme outil. Je n'avais jamais entendu cette expression auparavant. Je trouve cela intéressant dans la mesure où cela nous amène à réfléchir à la manière dont nous limitons certains de ces modèles à des domaines spécifiques et dont nous restreignons leur portée.

J'aimerais vous poser une question concernant votre connaissance d'autres organisations ou de la vôtre qui travaillent dans ce domaine. Y a-t-il quelqu'un qui se lance dans l'élaboration d'un code modèle, quelque chose que les pays du monde entier pourraient examiner alors que nous nous engageons dans cette voie pour tenter de réglementer ou d'établir très rapidement des paramètres de sécurité dans un secteur en pleine évolution? Comment des organismes tels que la vôtre aident-ils les parlementaires et les législateurs du monde entier à s'orienter vers une approche appropriée de cette question?

Anthony Aguirre: C'est une excellente question. Il existe une sorte de problème frustrant de l'œuf et de la poule, dans la mesure où il est difficile de renforcer les capacités de gouvernance en matière d'évaluations, de certifications et de toute l'infrastructure nécessaire pour évaluer et tester ces modèles en l'absence de toute forme d'obligation. Il ne peut pas y avoir de client sans qu'une forme de réglementation se mette en place.

D'un autre côté, quand on pense à la réglementation, cela semble très intimidant de voir qu'il existe tous ces différents cas d'utilisation pour ces systèmes, et qu'il faut réfléchir à la manière dont on va réglementer tout cela.

Comme M. Tegmark l'a souligné et comme vous l'avez suggéré, je pense qu'il est essentiel d'identifier les premières mesures à prendre. Il s'agit peut-être de tests de sécurité pour les enfants ou simplement d'exigences selon lesquelles, lorsque vous produisez un système d'IA orienté vers ce type d'outil, vous devez préciser à quoi il sert, puis vous pouvez commencer à évaluer si ce système d'IA est adapté à cet usage.

Le président: Je vous remercie, monsieur Aguirre. Nous allons à présent passer à un autre intervenant.

[Français]

Monsieur Thériault, vous avez la parole pour cinq minutes.

Luc Thériault: Merci, monsieur le président.

Monsieur Aguirre, je vous cite: « Les dirigeants de DeepMind, OpenAI et Anthropic [...] ont tous littéralement signé une déclaration affirmant que l'IA avancée pose un *risque d'extinction pour l'humanité*. »

Vous dites que c'est sans précédent, étant donné qu'ils construisent ces systèmes « sous des incitations commerciales et une surveillance gouvernementale quasi nulle ».

Que doit-on penser des entreprises qui font des mises en garde contre leurs propres produits, tout en continuant à les développer?

J'aimerais aussi entendre la réponse de M. Krueger.

[Traduction]

Anthony Aguirre: C'est assez étonnant. Nous n'avons jamais vu un secteur développer une nouvelle technologie tout en admettant publiquement sa dangerosité. C'est le résultat du développement particulier de ce secteur et, en particulier, de la course effrénée dans laquelle ces entreprises se sont lancées.

Il y a quelques semaines, à Davos, deux dirigeants d'entreprises spécialisées dans l'IA ont déclaré vouloir ralentir le rythme. Ils s'inquiètent de leurs activités, mais estiment ne pas pouvoir faire autrement, car ils sont engagés dans une course effrénée. S'ils lèvent le pied, ils risquent de finir perdants, car tous leurs concurrents ne se priveront pas d'appuyer sur l'accélérateur. Toutes ces entreprises es-

timent qu'elles doivent développer ce type de technologie, car quelqu'un d'autre le fera, et elles pensent que si quelqu'un doit le faire, autant que ce soit elles.

C'est une situation absurde pour nous, tout comme la course classique à l'armement qui a abouti à 70 000 ogives nucléaires, ce que personne ne trouvait pourtant excessif à l'époque. C'est là où nous en sommes arrivés à cause de la course à l'armement. Nous sommes ici dans une situation similaire, où il faut un acteur extérieur, et ce doit vraiment être le gouvernement, pour mettre fin à la course. Les grandes entreprises ne pourront pas y parvenir seules, même si elles le souhaitent.

• (1625)

David Krueger: Je suis d'accord avec les propos de M. Aguirre. Je tiens d'ailleurs à rappeler que l'idée d'une extinction de l'espèce humaine a été utilisée pour faire passer le message qu'il s'agit d'une sorte de coup de pub des entreprises. C'est manifestement faux. En tant que personne travaillant dans ce domaine depuis bien plus longtemps que ces entreprises n'existent, dans le cas d'Anthropic ou d'OpenAI, ces préoccupations remontent à plusieurs décennies.

Il est louable que beaucoup de chefs d'entreprise aient reconnu ces risques en même temps. Peut-être qu'ils exagèrent parfois les choses afin de promouvoir leur produit ou autre, mais je pense que c'est essentiellement ce qui se passe. Ils sont désespérés parce qu'ils se rendent compte qu'ils sont à quelques années de construire quelque chose qui leur fait peur, et ils veulent que quelque chose intervienne et désamorce la course, si possible.

Néanmoins, ces chefs d'entreprise ne croient pas que cela soit possible, en général, et je pense que c'est là leur erreur. Nous devons comprendre que c'est possible. C'est pourquoi j'ai mentionné les puces informatiques comme point d'intervention clé. Si aucune de ces entreprises et personne à l'échelle mondiale ne peut avoir accès à ces énormes quantités de puces et à l'énergie nécessaire pour construire ces centres de données, alors cette course ne se poursuivra pas au rythme actuel.

Par ailleurs, j'ai l'impression que la raison pour laquelle ces chefs d'entreprise agissent ainsi est qu'ils pensent que s'ils cessent d'innover en matière d'IA, leurs concurrents ne s'en priveront pas. Il y a le désir de gagner beaucoup d'argent. Certaines personnes souhaitent également voir apparaître cette nouvelle espèce dont parlait M. Tegmark. Elles veulent remplacer l'espèce humaine par l'IA, considérant qu'il s'agit là de la prochaine étape naturelle dans l'évolution. De nombreux acteurs influents au sein de l'industrie de l'IA ont fait des commentaires publics dans ce sens, et il s'agit là d'une vision du monde incroyablement misanthrope.

[Français]

Luc Thériault: J'aimerais amener la conversation sur une conclusion.

Monsieur Aguirre, vous appelez à un accord international entre les États-Unis, la Chine et d'autres pays qui peuvent se doter d'un solide mécanisme de vérification pour garantir que les partis et les rivaux ne feront pas de défection.

Monsieur Krueger, vous affirmez qu'arrêter la course à la construction de la superintelligence est raisonnable, possible et que ça constitue un impératif moral. Je suis assez d'accord là-dessus. Que peut faire le Comité pour lancer des négociations ou arriver à un tel accord?

Ce sujet revient. Vous n'êtes pas les seuls à dire ça. Que fait-on de cela?

Mes questions s'adressent à MM. Aguirre et Krueger.

[Traduction]

Le président: Monsieur Aguirre, il vous reste une dizaine de secondes, je vous prie.

Anthony Aguirre: Le Canada peut certainement adopter une position de protection de soi-même en affirmant que nous ne devrions pas développer de superintelligence artificielle. Je pense qu'il sera essentiel, pour mettre fin à cette course, que les États-Unis et la Chine réalisent qu'il est contraire à leurs propres intérêts de développer une superintelligence artificielle. Tant que les puissants de ce monde continueront de croire que le développement de l'IA peut leur conférer encore davantage de pouvoir, ils souhaiteront évidemment poursuivre dans cette voie. Mais en réalité, la superintelligence artificielle risque à terme d'absorber le pouvoir plutôt que de le conférer à tel ou tel être humain. S'ils en prennent conscience, il sera alors dans leur intérêt, ainsi que dans celui de tous de ne pas la développer, et c'est là le fondement même de tout l'enjeu.

Le président: Merci, monsieur Aguirre.

Monsieur Krueger, je vais vous laisser 10 secondes pour répondre.

David Krueger: Merci.

Nous devons entamer immédiatement des négociations pour conclure un traité international. Je veux que toutes les personnes ici présentes parlent à tous ceux qui ont la capacité de concrétiser une telle chose et leur transmettent le message suivant: dites-leur que vous avez entendu les témoignages de nous tous et d'experts précédents.

Le président: Merci, monsieur.

[Français]

Monsieur Hardy, vous avez la parole pour deux minutes et demie.

Gabriel Hardy: Merci beaucoup.

Je vais à nouveau poser ma question à M. Krueger, parce que c'est une question importante à analyser.

On voit beaucoup d'entreprises investir dans l'intelligence artificielle dans le but de remplacer leurs employés et de faire un peu plus de profits. J'ai soulevé cette question la semaine dernière, en comité, et elle a suscité de bonnes réponses. Les gens se questionnent.

Lorsque les grandes entreprises technologiques investissent dans l'intelligence artificielle, comment mesurent-elles le rendement de leur investissement? Voient-elles vraiment un bon rendement, ou est-ce que, finalement, en devant gérer l'intelligence artificielle et s'assurer que le travail qu'elle fait est correct, ça leur coûte plus cher que d'avoir un employé bien formé pour faire le même travail?

J'aimerais vous entendre à cet égard.

• (1630)

[Traduction]

David Krueger: Il y aura des cas où ce sera plus ou moins cher à l'heure actuelle. Il convient de reconnaître que l'IA n'a pas besoin de mieux accomplir une tâche si elle peut le faire à coût moindre. Des êtres humains compétents pourraient être remplacés par une IA

moins compétente, simplement parce que c'est moins cher, ce qui aurait des conséquences néfastes, à mon avis.

Par ailleurs, je pense que nous devons réfléchir à la direction que prend tout cela, car les choses évoluent très rapidement depuis quelques années. Je pense que d'ici quelques années, nous verrons une IA extrêmement compétitive qui remplacera la plupart des travailleurs. C'est sur cette prémisse que reposent ces investissements. Les investissements massifs pour développer de l'IA se justifient par la croyance que cela mènera, ou du moins pourrait mener, à la création d'une sorte de superintelligence telle que nous l'avons décrite ici.

[Français]

Gabriel Hardy: Les entreprises ne font jamais d'investissement sans regarder des indicateurs clés de rendement, ou ICP, et le taux de rendement. On parle toujours de l'avenir, mais aujourd'hui, est-ce qu'elles mesurent le taux de rendement? Est-ce qu'elles voient un rendement, ou diriez-vous que, pour l'instant, elles ne mesurent rien quant à leurs investissements dans cette nouvelle technologie?

[Traduction]

Le président: Répondez en 35 secondes maximum, je vous prie.

David Krueger: Je suis sûr que certains éléments sont mesurés. Depuis longtemps, les entreprises de la Silicon Valley se soucient davantage de la croissance, de la domination du marché et de la fidélisation des clients avant même de réaliser des profits. Amazon n'a pas généré de bénéfices pendant une dizaine d'années. Je ne sais pas ce que les entreprises recherchent actuellement, mais je ne crois pas que c'est ce qui stimule les investissements. Je pense que si l'on prend au sérieux la possibilité d'une intelligence artificielle générale, ou AGI, et d'une superintelligence dans quelques années, ce que font les investisseurs et surtout les personnes qui les développent, alors les investissements sont certainement justifiés, sauf que c'est aussi incroyablement dangereux et ne devrait pas se produire. Si cela tue tout le monde, cela tuera les investissements, les propriétaires de ces entreprises et toutes les personnes présentes dans cette salle. Le fait d'avoir créé une foule de bonnes entreprises n'aura pas d'importance, comme le dit Sam Altman. Peu importe que nous ayons guéri le cancer, etc.

Le président: Merci, monsieur Krueger. Je vous suis reconnaissant de ces observations.

[Français]

Monsieur Sari, vous avez la parole pour deux minutes et demie.

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

Je remercie les témoins de leurs interventions très pertinentes.

Je ne dispose que de deux minutes et demie pour poser des questions sur un sujet qui est très complexe et qui me tient vraiment à cœur.

Nous sommes tous d'accord sur le risque que vous avez expliqué aujourd'hui. Il s'agit de savoir comment intervenir. Pour ce faire, d'abord, il faut préciser le domaine ou le type d'intelligence artificielle dont on parle. Ici, on ne parle pas de l'intelligence artificielle générative, et on ne parle pas nécessairement de la superintelligence ou de l'intelligence artificielle agentive. La concentration du pouvoir cognitif entre les mains de l'intelligence artificielle est l'un des éléments qui m'inquiètent beaucoup, car c'est ce pouvoir cognitif qui oriente l'apprentissage d'une manière générale et la « cognitivité » universelle.

J'aimerais vraiment que l'un d'entre vous réponde à la question suivante: comment peut-on parler d'un contrôle humain ou gouvernemental réel, lorsque des systèmes d'intelligence artificielle apprennent, évoluent et prennent des décisions plus vite que nous pouvons, collectivement, les comprendre et les contester?

[Traduction]

Le président: Qui veut répondre à cette question en un peu plus d'une minute?

Monsieur Tegmark, puis-je commencer avec vous?

Max Tegmark: Oui.

Ce que vous décrivez avec tant d'éloquence, c'est la recherche numérique sur le gain de fonction, également connue sous le nom d'auto-amélioration récursive, où l'IA crée une meilleure IA qui crée à son tour une meilleure IA.

Encore une fois, c'est facile à gérer. Nous avons déjà traité ce problème en biologie ici, aux États-Unis. Nous avons interdit la recherche sur le gain de fonction, qui fait désormais l'objet d'une très vive opposition depuis qu'il est possible qu'elle ait causé la pandémie de COVID. Nous devrions également interdire la recherche sur le gain de fonction numérique dans le domaine de l'IA. Cela va de soi, mais un certain nombre d'entreprises américaines redoublent explicitement d'efforts dans ce type de recherche sur le gain de fonction numérique parce qu'elles ne sont pas réglementées.

• (1635)

Le président: Merci, monsieur Tegmark.

Messieurs Krueger ou Aguirre, avez-vous quelque chose à ajouter en 15 ou 20 secondes?

David Krueger: Oui.

Je pense qu'il est important de souligner que ce problème doit être traité au niveau international. C'est aussi quelque chose que M. Aguirre et M. Tegmark ont évoqué en parlant de l'IA... Je pense que c'est un objectif louable, mais nous devons peut-être prendre des mesures draconiennes dans un proche avenir pour pouvoir surveiller et faire respecter un accord international visant à mettre fin à cette course. Ensuite, lorsque nous aurons la situation sous contrôle en quelque sorte, nous pourrions réfléchir à la façon dont nous pouvons procéder pour développer une IA utile...

Le président: Merci, monsieur Krueger. Je dois vous interrompre. C'est la fin de l'heure.

Au nom du Comité, je tiens à vous remercier tous les trois d'avoir participé à cette discussion. Merci.

Nous allons suspendre la séance quelques minutes pendant que nous passons à la deuxième heure de la réunion avec le commissaire à la protection de la vie privée.

• (1635)

(Pause)

• (1640)

Le président: Nous allons reprendre la séance pour la deuxième heure afin de poursuivre notre étude des défis que posent l'intelligence artificielle et son encadrement.

Je veux souhaiter la bienvenue pour la deuxième heure de réunion aujourd'hui à M. Philippe Dufresne, commissaire à la protection de la vie privée du Canada — c'est toujours un plaisir de vous accueillir parmi nous, monsieur —, et à Marc Chénier, sous-commissaire et avocat général principal, des Commissariats à l'information et à la protection de la vie privée au Canada.

Monsieur Dufresne, vous avez au plus cinq minutes pour faire votre déclaration. On vous écoute, je vous prie.

[Français]

Philippe Dufresne (commissaire à la protection de la vie privée du Canada, Commissariats à l'information et à la protection de la vie privée au Canada): Merci beaucoup, monsieur le président.

Mesdames et messieurs les membres du Comité, je vous remercie de cette invitation à comparaître dans le cadre de votre étude sur les défis que posent l'intelligence artificielle et son encadrement.

Tenir compte des répercussions sur la vie privée des progrès technologiques qui se succèdent à un rythme rapide est l'une de mes priorités stratégiques. Cela a également été au centre de mes travaux nationaux, internationaux et interréglementaires au cours des dernières années, compte tenu de son adoption rapide et à grande échelle par les individus et les organisations au Canada et dans le monde entier.

[Traduction]

La protection de la vie privée est une question importante et d'actualité au Canada. Alors que de plus en plus de données personnelles sont recueillies, utilisées et communiquées, la protection des données revêt une importance croissante pour la population et les organisations canadiennes.

La protection des renseignements personnels est particulièrement importante dans le contexte de l'IA, car ces renseignements peuvent être utilisés pour entraîner et exploiter ces systèmes. Récemment, j'ai annoncé la tenue d'une enquête approfondie sur la plateforme de média social X et son robot conversationnel Grok. L'enquête portera sur le phénomène émergent de l'utilisation de l'IA pour créer des hypertrucages, qui peuvent présenter un risque important pour la population canadienne, y compris les enfants.

[Français]

Je m'attends à ce que les résultats de cette enquête, ainsi que mon enquête en cours sur OpenAI, contribuent à orienter les politiques et les mesures de protection de la vie privée en ce qui a trait à l'intelligence artificielle et aident les individus et les organisations à utiliser et à déployer ces technologies de manière sûre et responsable et en bénéficiant de protections adéquates pour les données personnelles.

Les enquêtes menées par le Commissariat à la protection de la vie privée au cours des deux dernières années ont démontré comment les lois canadiennes permettent de traiter les enjeux importants liés à la protection de la vie privée qui peuvent avoir de graves répercussions sur les individus.

[Traduction]

Par exemple, mon enquête sur Aylo, qui exploite Pornhub et d'autres sites Web pornographiques, portait sur le partage non consensuel d'images intimes. Mon enquête conjointe avec mon homologue britannique sur l'atteinte survenue à 23andMe portait sur un incident qui a compromis les renseignements personnels très sensibles de 7 millions de clients, dont plus de 300 000 Canadiennes et Canadiens.

L'automne dernier, j'ai annoncé les résultats de mon enquête menée avec mes homologues provinciaux sur TikTok, qui a mis en évidence l'importance de protéger la vie privée des enfants en ligne. À la suite de notre enquête, l'entreprise a apporté des améliorations à ses pratiques en matière de protection de la vie privée et continue de le faire dans l'intérêt supérieur de ses utilisateurs, en particulier des enfants.

Les technologies telles que l'IA peuvent apporter des avantages économiques, sociaux et d'intérêt public. La valeur de cette innovation sera maximisée lorsqu'elle s'accompagnera de confiance.

[Français]

Un sondage mené par le Commissariat l'année dernière a révélé qu'une grande majorité de Canadiennes et de Canadiens s'inquiètent de la manière dont leurs renseignements personnels sont recueillis et utilisés. En effet, 83 % d'entre eux se disent préoccupés par la protection de leur vie privée lorsqu'ils utilisent des outils d'intelligence artificielle. Plusieurs personnes ont pris des mesures pour se protéger, et la plupart ont indiqué être moins disposées qu'il y a cinq ans à communiquer leurs renseignements personnels à des organisations.

[Traduction]

Cela met encore davantage en évidence l'avantage stratégique pour les organisations de développer et de déployer l'IA, ainsi que d'autres technologies, de manière responsable et dans le respect de la vie privée. Il est important pour les développeurs et les fournisseurs d'IA d'intégrer la protection de la vie privée dans la conception, l'exploitation et la gestion des nouveaux produits et services. Nous leur demandons aussi de tenir compte de l'incidence particulière que ces outils ont sur les enfants ainsi que sur les groupes qui sont depuis longtemps victimes de discrimination ou de préjugés.

Les organisations qui utilisent l'IA devraient faire preuve de transparence et être tenues responsables de toute décision générée par l'IA à propos d'individus, qu'il s'agisse d'accorder un prêt ou de donner un emploi à quelqu'un.

• (1645)

[Français]

Alors que les technologies continuent d'évoluer rapidement et s'intègrent de plus en plus dans notre vie personnelle et professionnelle, il est de notre devoir collectif, en tant qu'organismes de réglementation et décideurs politiques, de veiller à la protection de la vie privée des générations actuelles et futures. Les lois canadiennes sur la protection des renseignements personnels doivent permettre de relever ce défi et, pour y arriver, ces lois doivent être modernisées.

[Traduction]

En ce qui concerne l'IA, les modifications que je recommande d'apporter aux lois fédérales canadiennes sur la protection des renseignements personnels comprennent la reconnaissance du droit à la vie privée comme un droit fondamental, ainsi que l'établissement d'exigences visant à prévoir des mesures de protection de la vie privée dès la conception et à réaliser des évaluations des facteurs relatifs à la vie privée pour les activités de traitement de données à incidence élevée.

Les renseignements personnels sont au cœur de l'intelligence artificielle, et c'est pourquoi les lois sur la protection des renseignements personnels devraient, selon moi, être au cœur de la réglementation de l'IA.

Merci. Je serai heureux de répondre à vos questions.

Le président: Merci, monsieur Dufresne.

Nous allons commencer avec M. Barrett, pour six minutes. Je serai strict avec le temps.

Monsieur Barrett, allez-y.

Michael Barrett: Parlez-moi de ce qu'on appelle les voitures espionnes chinoises. Le premier ministre de l'Ontario, M. Ford, a exprimé de sérieuses préoccupations au sujet du plan du marché canadien d'accepter près de 50 000 véhicules fabriqués par des entreprises telles que BYD. Qu'est-ce que votre commissariat a examiné jusqu'à présent au sujet de ces véhicules et de cette affirmation, fondée ou non, de M. Ford?

Philippe Dufresne: J'ai entendu ces affirmations, et nous surveillons la situation de façon générale concernant les véhicules connectés. En fait, nous avons lancé cette année notre programme de contribution sur le thème des appareils connectés, et nous sommes impatients d'en savoir plus sur les types de connexion, les types de données recueillies par les véhicules et d'autres types d'appareils.

Nous ne nous penchons pas sur cette question précisément sous l'angle chinois. Toutefois, dans le cadre de notre enquête sur TikTok, l'un des éléments que nous avons soulignés dans notre conclusion était que lorsque des données quittent le Canada et qu'il existe un risque que d'autres gouvernements puissent y avoir accès, les Canadiens devraient en être informés. Ce doit être transparent. Dans les conclusions de notre rapport sur TikTok, nous avons demandé de l'énoncer explicitement, et l'organisation a accepté.

Michael Barrett: En ce qui concerne TikTok, avez-vous une recommandation à faire aux Canadiens quant à savoir s'ils devraient utiliser l'application ou non?

Philippe Dufresne: Nous recommandons aux Canadiens de poser des questions sur cette application ou n'importe quelle application. Pour être franc, dans toute situation où leurs renseignements personnels sont demandés, ils devraient poser les questions suivantes: « Pourquoi en avez-vous besoin? Qu'en ferez-vous? À qui les communiquerez-vous? » Dans le cadre de l'enquête sur TikTok, nous abordons ces questions de front et examinons ce que l'organisation dit aux Canadiens.

C'est une chose, que les citoyens posent des questions. Les organisations ont la grande responsabilité d'être proactives dans cette transparence. Dans le cas de TikTok, nous avons constaté que l'information n'était pas suffisamment claire pour les adultes, et elle ne l'était certainement pas pour les enfants, qui représentent une part importante de ce marché.

Les questions devraient porter sur l'utilisation, la diffusion, les fins auxquelles les données sont utilisées, la destination et les personnes qui y ont accès. Les Canadiens devraient poser davantage de questions de ce genre, mais les organisations devraient assumer proactivement leurs responsabilités en rendant ces informations facilement accessibles.

Michael Barrett: Avez-vous constaté un changement dans la communication proactive d'informations par TikTok aux Canadiens depuis que vous avez formulé ces recommandations?

Philippe Dufresne: Je dirais oui, car nous travaillons avec ces organisations pour surveiller la mise en œuvre de nos recommandations. Nous avons six mois pour les mettre en place. Une recommandation importante était d'améliorer les outils pour empêcher les enfants mineurs d'accéder à la plateforme. D'autres portaient sur la transparence, le consentement et l'information.

Un certain nombre de recommandations ont été mises en œuvre, et les organisations ont jusqu'au mois de mars pour terminer le reste. Nous allons surveiller cela pour veiller à ce que ce soit fait.

Michael Barrett: Tout d'abord, compte tenu de leur coopération jusqu'à présent, pensez-vous qu'elles respecteront le délai? Si elles ne le respectent pas, quelles seront les conséquences?

● (1650)

Philippe Dufresne: Si elles ne respectent pas le délai, l'une des lacunes dans la loi sur la protection des renseignements personnels est que je n'ai pas le pouvoir de rendre des ordonnances. Nous devons compter sur leur collaboration et travailler avec elles. Jusqu'à présent, cette collaboration fonctionne et j'ose espérer qu'elle se poursuivra, mais c'est pourquoi un élément clé de la réforme de la loi doit être d'accorder le pouvoir de rendre des ordonnances au commissaire à la protection de la vie privée.

Michael Barrett: Je n'ai plus beaucoup de temps, alors je vais passer à un autre sujet.

Je veux vous interroger au sujet de votre enquête sur X Corp. C'est évidemment extrêmement important, en raison des dommages alarmants que causent en temps réel, non seulement aux enfants, mais surtout aux enfants, la propagation et la création d'images non consenties.

Dans votre rôle de commissaire canadien, quelle est votre position étant donné que vous n'avez pas le pouvoir de rendre des ordonnances? Manquez-vous d'outils d'application nécessaires pour vous acquitter efficacement de votre rôle pour protéger les Canadiens, surtout les enfants, dans une situation dont la gravité n'échappe à personne?

Philippe Dufresne: Je dirai deux choses.

La première est que je dois détenir le pouvoir de rendre des ordonnances. Cela devrait être une priorité absolue pour le Parlement. Les déclarations du ministre de l'IA selon lesquelles une modernisation est envisagée m'encouragent.

Cela dit, j'utilise et je continuerai d'utiliser les outils existants dont je dispose actuellement pour protéger la vie privée des Cana-

diens en menant des enquêtes, en formulant des recommandations et en collaborant avec des partenaires au Canada et à l'étranger.

En ce qui concerne les hypertrucages, nous avons publié des déclarations dès décembre 2023 pour les dénoncer. Nous travaillons avec des partenaires internationaux pour mobiliser la communauté internationale afin qu'elle prenne des mesures.

C'est une chose que j'ai faite, par exemple, dans l'affaire 23andMe concernant la violation des données des Canadiens. Je n'ai pas le pouvoir de rendre des ordonnances, mais mon homologue britannique l'a, et ensemble, nous avons agi dans l'intérêt des Canadiens.

Le président: Merci, messieurs Dufresne et Barrett.

[Français]

Monsieur Sari, vous avez la parole pour six minutes.

Abdelhaq Sari: Merci, monsieur le président.

Je vous remercie beaucoup de votre témoignage, monsieur Dufresne.

Avant de poser ma question, j'aimerais bien que nous revenions un peu sur votre mission quant à la protection de la vie privée de manière générale. Celle-ci repose sur des principes fondamentaux, notamment sur la question de la transparence, la responsabilisation de manière générale et la capacité pour un citoyen ou une citoyenne de pouvoir comprendre, voire contester les décisions qui sont prises par n'importe quelle solution, que ce soit une solution d'intelligence artificielle ou une autre.

Cependant, vous savez très bien que l'intelligence artificielle transforme un peu la manière dont les décisions sont prises. Par exemple, qu'elle soit générative, cognitive ou agentique, n'importe quelle intelligence artificielle se passe beaucoup sur le réseau de neurones, qui ne donne pas ce processus décisionnel. C'est ce qu'on appelle la boîte noire. Il n'y a pas de modèle qui est donné sur la décision.

Ma première question est la suivante.

Dans un contexte comme celui-ci, où les technologies évoluent de manière très rapide, comment le Commissariat évolue-t-il? Comment évalue-t-il l'existence d'un contrôle humain réel? Est-ce qu'on peut parler d'un certain contrôle humain?

Philippe Dufresne: Je vous remercie de votre question.

Vous avez tout à fait raison. Les principes de protection de la vie privée sont des principes qui peuvent évoluer dans le temps, parce que ce sont des principes fondamentaux qui touchent la dignité humaine, la liberté, la capacité de prendre des décisions, la capacité de choisir ce qu'on communique et à qui on le communique. Maintenant, on fait face à des défis importants avec la technologie, parce que celle-ci évolue très rapidement. Il faut donc que des organismes comme le mien et d'autres, en général, se rattrapent. Il faut qu'on développe une connaissance de ces outils, qu'on développe la capacité technologique. Nous avons créé, à mon bureau, cette capacité de comprendre la technologie et de pouvoir agir en conséquence.

Pour ce qui est de la question de la décision humaine à laquelle vous faites allusion, il faut aussi qu'on soulève ces enjeux dans le cadre de réformes législatives. J'ai recommandé, par exemple, qu'on ait une transparence beaucoup plus grande dans le contexte de l'intelligence artificielle pour comprendre ce qui a été décidé et qu'on s'assure qu'il y a un humain dans le processus. Tout récemment, j'étais à Séoul pour l'Assemblée mondiale pour la protection de la vie privée et on a mis en avant, à mon commissariat, une résolution qui a été adoptée par la communauté internationale. Elle portait justement sur le rôle de l'humain dans le processus de l'intelligence artificielle, son importance et ce qu'il devrait représenter. Cet enjeu est donc très pertinent et ça vient contrer un peu l'effet de boîte noire que vous avez mentionné.

• (1655)

Abdelhaq Sari: La boîte noire est plus technologique, c'est la manière dont fonctionnent les réseaux de neurones d'une manière générale et on ne peut pas revenir là-dessus pour la changer. Ça date des années 1980 et ça va continuer de fonctionner de la sorte.

Vous parlez de comprendre la technologie; estimez-vous que le droit à la protection de la vie privée est toujours pleinement respecté, présentement?

Philippe Dufresne: Je pense qu'il y a certainement des cas où ce n'est pas respecté. Nous avons parlé de la décision concernant TikTok, où il y avait eu des manquements, ou de la décision concernant Bureau en gros qui a été rendue récemment. Il y a beaucoup de cas où il y a des manquements, mais ce qui est important, c'est d'avoir les outils pour pouvoir agir contre ces ceux-ci. Idéalement, on les prévient avant qu'ils ne se produisent, sinon, on agit après le fait.

À mon avis, ces principes sont neutres sur le plan technologique et ils peuvent évoluer avec la technologie. Cependant, la législation doit être modernisée pour rendre ça plus facile. Des concepts comme la désidentification, la transparence, la façon de traiter la transparence dans une situation où les organismes vont dire qu'ils ne savent pas eux-mêmes comment ils arrivent à pareilles décisions. Le président de Google a dit ça récemment. La législation doit donc évoluer. Nous pouvons avoir des obligations plus proactives, mais nous avons déjà un très bon cadre en ce sens avec la Loi sur la protection des renseignements personnels.

Abdelhaq Sari: Ça me rassure de vous entendre parler de cadre, parce que tous les articles scientifiques que je lis, ainsi que les gens qui sont venus témoigner aujourd'hui, parlent beaucoup plus de l'opacité des technologies d'intelligence artificielle. Je comprends votre rôle et je le respecte énormément. Je trouve qu'il est extrêmement important. Vous me rassurez pas mal, aujourd'hui, compte tenu des avancées que vous avez faites dans votre travail.

De notre côté, en tant que comité de la Chambre des communes, que nous recommandez-vous de faire sur le plan du contrôle? Qu'est-ce que nous devons faire pour orienter cette capacité de contrôle, sachant qu'il y a quand même une certaine opacité, qu'on le veuille ou non?

Philippe Dufresne: À mon avis, la chose la plus importante que vous pouvez faire comme parlementaires, c'est modifier la Loi. C'est un outil essentiel. Moi, j'ai fait sept recommandations prioritaires au gouvernement, je les ai envoyées au Comité. Ce sont des éléments comme des pouvoirs accrus d'exécution de la loi, le pouvoir d'émettre des ordonnances, le pouvoir d'imposer des amendes dans certains cas. La reconnaissance du droit à la vie privée comme droit fondamental est importante. C'est important plus que jamais

avec l'arrivée de l'intelligence artificielle. Il s'agit d'éléments comme la reconnaissance de la vie privée dès la conception des outils. Il est important de demander aux organisations qui vont créer des outils ayant un impact important de faire des évaluations des effets sur la vie privée au début, pas à la fin. Ce serait positif sur tous les plans. Il faut aussi traiter la question sur le plan international. Il faut avoir un cadre régissant les données qui quittent le Canada.

Abdelhaq Sari: En matière d'information, je dis toujours qu'il faut colliger l'information, la traiter, la stocker, la sécuriser et pouvoir la transmettre à l'échelle internationale avec d'autres facteurs. C'est dans la sécurisation et la protection que votre rôle prend tout son sens.

Revenons aux citoyens qui nous écoutent aujourd'hui. Y a-t-il une catégorie de citoyens et de citoyennes plus vulnérable? Qui sont les plus vulnérables? On a parlé beaucoup des jeunes et des femmes, aujourd'hui. Y a-t-il des gens plus vulnérables que d'autres?

Le président: La réponse devra être très brève, monsieur Dufresne.

Philippe Dufresne: Les jeunes sont très vulnérables, parce qu'ils sont assujettis à cela, ils vivent là-dedans sur tous les plans. Vous avez raison de mentionner les femmes et les personnes âgées. Je pense que beaucoup de groupes ont besoin de protection.

Le président: Merci, monsieur Sari.

Monsieur Thériault, vous avez la parole pour six minutes.

Luc Thériault: Merci, monsieur le président.

Soyez le bienvenu, monsieur le commissaire.

Vous avez parlé tout à l'heure du volet international. Je voulais vous ramener...

Abdelhaq Sari: Monsieur le président, je n'entends pas bien M. Thériault.

Le président: Attendez un instant.

Veuillez abaisser votre micro, monsieur Thériault. C'est mieux comme ça. Merci.

Je redémarre le chronomètre.

Luc Thériault: D'accord, merci.

Je voulais vous emmener à la déclaration des dirigeants du G7 qui dit ceci:

Travailler ensemble afin d'accélérer l'adoption de l'IA dans le secteur public pour améliorer la qualité des services publics offerts aux citoyens et aux entreprises et accroître l'efficacité du gouvernement, tout en respectant les droits de la personne et la vie privée et en faisant la promotion de la transparence, de l'équité et de la responsabilité.

Un peu plus loin, on y dit ceci:

Promouvoir la prospérité économique en aidant les PME à adopter et à développer une IA qui respecte les données personnelles et les droits de propriété intellectuelle et renforce leur état de préparation, leur efficacité, leur productivité et leur compétitivité.

Pour les dirigeants du G7, quels devraient être les éléments clés d'une approche centrée sur l'être humain à l'égard de l'intelligence artificielle, d'après vous?

• (1700)

Philippe Dufresne: Je vous remercie de votre question.

Vous avez fait allusion à la déclaration des dirigeants du G7. Nous, les commissaires à la vie privée du G7, avons tenu notre propre réunion ici, au lac Meech. Nous avons émis, en juin, une déclaration unanime sur l'importance, justement, d'accorder la priorité à la protection de la vie privée pour deux raisons: favoriser l'innovation et l'économie et protéger les enfants et les jeunes.

Nous avons fait allusion aux déclarations des dirigeants sur ces éléments, et nous avons donné des exemples concrets: tenir compte de l'intérêt supérieur des enfants et des jeunes quand on conçoit des technologies, quand on les développe et quand on les met sur le marché; souligner l'importance de commencer la protection de la vie privée dès le début, et non à la fin; souligner l'importance de la transparence et de la communication avec les citoyens pour générer de la confiance.

C'est pour cette raison que certains disent parfois qu'un régime de protection de la vie privée robuste peut freiner l'innovation. Je le conteste, et notre déclaration du G7 le contestait aussi. Elle dit très clairement, en fait, que ça va encourager l'innovation et soutenir le développement économique, parce que ça va générer la confiance des citoyens et celle des consommateurs devant cette nouvelle technologie. On le voit dans les sondages auxquels j'ai fait allusion un peu plus tôt: il y a des préoccupations à ce sujet. C'est donc essentiel du point de vue des droits fondamentaux, si on veut protéger à la fois les enfants, les femmes, les aînés, nous tous et nous toutes, pour ce qui est de notre liberté et de notre dignité. De plus, de façon pragmatique, ça va soutenir l'économie mondiale et l'économie canadienne en particulier. C'est un message fort qui a été émis ici, à Ottawa. J'étais très fier de cette déclaration, et nous continuons d'en faire la promotion, ici et ailleurs.

Luc Thériault: En ce qui concerne l'amélioration de la qualité des services publics et de l'accroissement de l'efficacité du gouvernement, comment peut-on concilier l'accélération de l'adoption de l'intelligence artificielle avec le respect des droits de la personne et de la vie privée et la promotion de la transparence, de l'équité et de la responsabilité?

Philippe Dufresne: Justement, il faut avoir cet équilibre. Les outils vont apporter une productivité accrue. Ils peuvent soutenir toutes sortes de développements économiques. Ils peuvent soutenir des développements sur le plan de la santé. Ils peuvent améliorer les services aux citoyens. C'est donc très positif. Toutefois, il faut le faire de façon transparente et de façon humaine. Il faut que la notion soit centrée sur l'humain. Il faut le faire de façon à respecter les principes fondamentaux de protection de la vie privée, par exemple le consentement éclairé, dans certains cas, ou les objectifs appropriés.

Au mois de septembre 2025, nous avons émis une analyse avec nos collègues du Bureau de la concurrence, ceux du Conseil de la radiodiffusion et des télécommunications, ou CRTC, et ceux de la Commission du droit d'auteur au sujet des médias synthétiques, ce qu'on appelle parfois l'hypertrucage. Nous avons fait référence aux enjeux à cet égard. Dans certains cas, la technologie peut aider à la création d'images. Toutefois, si c'est utilisé pour manipuler les gens ou si c'est utilisé pour créer des images sexuelles sans le consentement des gens, bien sûr, ça viole des droits qu'il ne faut pas violer. Il ne faut pas mettre l'innovation d'un côté et la protection de la vie privée de l'autre. Les Canadiens et les Canadiennes peuvent et doivent avoir les deux.

Luc Thériault: Prenons l'exemple du travail. C'est un droit fondamental. On pourrait dire que ça fait partie des droits de la per-

sonne. Des gens sont venus nous dire que les États voudront bientôt remplacer les cols blancs par l'intelligence artificielle. À partir du moment où une masse critique d'individus perdent leur travail, comment concilier ça avec cette avancée de l'intelligence artificielle? On parle ici du G7, dont les participants ont donné des commandes de productivité et d'efficacité à l'État.

Philippe Dufresne: Je pense que toute la question relative à la protection des emplois outrepassa un peu de mon mandat de protection de la vie privée. Cependant, je peux vous dire que, quand on protège la vie privée, on protège les individus, on rend les éléments plus transparents, on permet de mieux comprendre ce qui se passe et on permet aussi de distinguer ce qui est humain et ce qui ne l'est pas. Je crois que ça peut faire partie des objectifs et des solutions. Dans le cadre des emplois, on pourrait dire qu'on ne veut peut-être pas concurrencer l'intelligence artificielle là où il s'agit purement de statistiques probabilistes, mais qu'il faut garder l'avantage humain. Selon moi, ça sera toujours l'avantage essentiel.

● (1705)

Le président: Merci, monsieur Dufresne et monsieur Thériault.

Nous allons commencer le deuxième tour de questions. M. Cooper dispose de cinq minutes.

[Traduction]

Allez-y, je vous prie, monsieur Cooper.

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

Commissaire Dufresne, le premier ministre a récemment conclu un accord avec Xi Jinping pour autoriser l'importation de 50 000 véhicules électriques chinois par année. Il l'a fait malgré les sérieuses préoccupations concernant la sécurité nationale et la protection des renseignements personnels, y compris la surveillance audio et vidéo et le suivi des déplacements des gens. Le premier ministre Ford a qualifié ces VE de véhicules espions. D'autres ont fait de même. Les États-Unis ont interdit les véhicules équipés de logiciels chinois.

Compte tenu de ces très graves préoccupations qui touchent directement la vie privée des Canadiens, Le Cabinet du premier ministre, le ministre du Commerce international ou un membre du gouvernement a-t-il consulté votre commissariat à ce sujet avant de conclure un accord avec Xi Jinping?

Philippe Dufresne: On ne nous a pas consultés à propos de cette initiative commerciale avec la Chine.

Michael Cooper: Merci.

Il y a environ une heure à la Chambre, le ministre responsable de l'IA a fait savoir que le gouvernement déposera une version mise à jour de la LPRPDE dans un proche avenir.

Vous a-t-on consultés à propos de cette version mise à jour?

Philippe Dufresne: Oui. Mon équipe et moi sommes en contact avec le ministère. Je suis ravi de la façon dont ces échanges se sont déroulés.

Michael Cooper: Vous attendez-vous à ce que la version mise à jour de la LPRPDE vous octroie le pouvoir de rendre des ordonnances dont vous avez besoin et que vous n'avez pas, comme vous l'avez souligné il y a quelques instants?

Philippe Dufresne: La décision reviendra au gouvernement qui présentera le projet de loi. J'espère et je pense avoir très clairement fait savoir au gouvernement que c'est l'une de mes principales priorités. J'ai besoin de détenir des pouvoirs d'application accrus.

Michael Cooper: D'accord.

Pour changer un peu de sujet, la Cour suprême, dans sa jurisprudence dans l'arrêt *Spencer* et l'arrêt *Bykovets*, a clairement déclaré que l'accès aux informations ou aux données personnelles exige une autorisation judiciaire préalable. Est-ce exact?

Philippe Dufresne: Dans les affaires criminelles... Dans les affaires de droit administratif, ce ne serait pas forcément la même chose.

Michael Cooper: En vertu du soi-disant projet de loi sur la cybersécurité du gouvernement libéral, le ministre de l'Industrie a le pouvoir d'ordonner à un fournisseur de services de télécommunication de recueillir, conserver et lui communiquer des informations telles que des métadonnées, des renseignements sur les comptes des abonnés, les visites de sites Web, l'emplacement des données financières, entre autres informations personnelles. Est-ce exact?

Philippe Dufresne: Oui.

Michael Cooper: La jurisprudence constitutionnelle a reconnu que les métadonnées telles que les activités de navigation sur Internet sont des renseignements personnels de nature très délicate qui sont protégées par l'article 8 de la Charte qui stipule que « toute personne a le droit d'être protégée contre les perquisitions et les saisies abusives ». Est-ce exact?

Philippe Dufresne: C'est ce que je comprends de cette jurisprudence. En vertu de mon projet de loi, les renseignements de nature délicate bénéficient d'un niveau de protection plus élevé.

Michael Cooper: C'est parce que les métadonnées peuvent révéler beaucoup d'informations sur les activités en ligne et la vie privée d'une personne, entre autres questions liées à la protection des renseignements personnels. Est-ce exact?

Philippe Dufresne: Oui. La Cour suprême, dans l'affaire *Bykovets*, a évoqué tout ce que l'on peut tirer de ce type d'information.

Michael Cooper: Le projet de loi C-8 n'exige pas du ministre de l'Industrie qu'il obtienne une autorisation judiciaire préalable comme condition pour ordonner la collecte, la conservation et la communication de renseignements personnels très confidentiels tels que les métadonnées, n'est-ce pas?

Philippe Dufresne: Non. Encore une fois, ce n'est pas dans un contexte criminel.

Je tiens à dire que j'ai comparu devant le Comité de la sécurité publique au sujet du projet de loi C-8 et que j'ai fait part de mes recommandations d'amendements au projet de loi, ce qui comprend la nécessité et la proportionnalité dans l'exercice de ces pouvoirs.

• (1710)

Michael Cooper: Croyez-vous que le ministre devrait être tenu d'obtenir une autorisation judiciaire? Devrait-on restreindre cette disposition, contrairement à ce qui est prévu dans le libellé actuel, où le ministre disposerait de vastes pouvoirs pour ordonner la collecte, la conservation et le transfert de renseignements personnels, ce qui, selon la Cour suprême, nécessite clairement une autorisation judiciaire?

Philippe Dufresne: Dans mon témoignage, j'ai souligné qu'il existe un équilibre entre les objectifs de sécurité nationale. C'est là où la nécessité et la proportionnalité deviennent importantes. Dans

certaines parties du projet de loi, j'ai souligné qu'il n'y avait que la nécessité ou...

Michael Cooper: Il serait juste de dire que cela ne suffit pas pour protéger les renseignements personnels des Canadiens. Il y a place à l'amélioration.

Philippe Dufresne: J'ai fait des recommandations sur la norme, sur les ententes d'échange de renseignements ainsi que sur la notification à mon commissariat en cas de violation.

Le président: C'est tout, monsieur Cooper.

Madame Church, la parole est à vous pour cinq minutes.

Leslie Church: Bienvenue, monsieur Dufresne. Je vous remercie encore une fois de votre présence au Comité et de tout le travail que votre commissariat accomplit pour faire la lumière sur certains de ces enjeux émergents et préoccupants en matière de protection de la vie privée, notamment en ce qui concerne l'IA et les outils numériques.

Je veux vous orienter dans une direction différente et plus positive. Dans une grande partie du travail que je fais auprès des aînés, des personnes handicapées, des familles et des aidants naturels, j'entends beaucoup parler des défis pour s'y retrouver dans les systèmes gouvernementaux, entre les systèmes provinciaux et fédéraux, et les processus de demande. L'une des choses qui m'intéressent dans le cadre d'un effort de modernisation est la façon dont nous pouvons mieux utiliser les données pour soutenir des modèles de service modernes, l'admissibilité proactive aux prestations et la simplification de l'accès aux services.

Y a-t-il quelque chose que vous proposez dans cette modernisation qui nous aiderait à mieux fournir des services en tant que gouvernement?

Philippe Dufresne: Je pense que c'est un point très important, car la protection de la vie privée ne doit pas et ne devrait pas faire obstacle à l'intérêt public. Les données peuvent contribuer à une bonne prestation des services publics. Elles peuvent même protéger la vie privée grâce aux technologies qui renforcent la confidentialité.

Dans une partie de mes recommandations, je parle de la protection de la vie privée dès la conception, avec une évaluation des facteurs relatifs à la vie privée, par exemple. C'est un excellent outil lorsque vous avez un processus dans lequel vous communiquerez des renseignements parce que vous devez le faire ou parce que vous allez les utiliser pour prendre de meilleures décisions pour les Canadiens. La transparence est importante. Cette évaluation est importante.

Dans bien des cas, lorsqu'on me demande mon avis, je ne dirai pas aux gens de ne pas faire quelque chose ou de ne pas donner à la police un outil s'il est justifié et nécessaire de le faire. Je vais demander si c'est nécessaire et proportionnel. Est-ce transparent? Avez-vous évalué la situation?

Mes recommandations de modifications tant à la loi sur la protection des renseignements personnels dans le secteur privé qu'à celle dans le secteur public vont dans ce sens. Par exemple, la loi sur la protection des renseignements personnels dans le secteur public n'exige pas une évaluation de la nécessité et de la proportionnalité. Je pense qu'elle devrait le faire. Je pense que cela permettra une communication et une utilisation accrues de l'information, mais d'une manière qui préserve la confiance des Canadiens.

Leslie Church: Cela protège les données et la vie privée. J'approuve cette idée.

Puis-je vous interroger sur le travail que votre commissariat a entrepris concernant le code de confidentialité des enfants? Pouvez-vous dire au Comité où vous en êtes à ce sujet et quelles seront les prochaines étapes?

Philippe Dufresne: Merci de la question.

Une réforme législative pourrait aussi aider dans ce domaine. En effet, la mouture précédente du projet de loi, le C-27 déposé lors de la dernière législature, conférerait le pouvoir d'adopter un code de pratique qui aurait été approuvé par moi et qui aurait instauré des garanties pour les organisations, de même que des obligations et de la prévisibilité.

Un code pour les enfants cadrerait très bien dans ce régime. Puisque ce cadre n'existe pas encore, le code pour les enfants que je m'efforce de mettre en place établira à l'intention des organisations mes attentes et les attentes du commissariat concernant des choses comme l'intérêt supérieur de l'enfant, l'activation systématique de mécanismes de protection des renseignements pour les contenus destinés aux enfants de même que l'emploi d'un langage adapté aux enfants pour la communication d'informations et de la politique sur le consentement.

Nos travaux qui se poursuivent comprennent des consultations avec des parties prenantes au Canada et partout dans le monde. Parmi les notions examinées, il y a l'estimation et la vérification de l'âge et les moyens pour empêcher les enfants qui n'ont pas l'âge minimal d'accéder à certains sites Web comme TikTok. Comment appliquer cela sans aller trop loin et sans recueillir une trop grande quantité de renseignements personnels? C'est à l'atteinte de cet équilibre que le commissariat tient à contribuer.

• (1715)

Leslie Church: Nous en sommes seulement aux balbutiements, mais vous inspirez-vous des enseignements tirés de la mise en oeuvre de la vérification de l'âge en Australie?

Philippe Dufresne: Oui. Nous en tenons compte. Nous collaborons très étroitement avec les collègues de l'Australie, du Royaume-Uni et de plusieurs autres pays. Nous avons publié des déclarations internationales sur ce que nous souhaiterions. Nous aimerions voir des mécanismes de réduction au minimum des données. Nous voulons quelque chose de transparent qui n'utilise pas les données à des fins autres que celles auxquelles elles ont été recueillies et qui protège contre les atteintes à la vie privée. Ce sont des questions sur lesquelles se penche vraiment la communauté internationale. L'Australie a dû bien évidemment accélérer le rythme en raison des mesures qu'elle a prises concernant les médias sociaux.

Leslie Church: Vous avez dit que le pouvoir de rendre des ordonnances était un des outils que vous aimeriez obtenir. Pourriez-vous nous dire quels autres outils ou capacités aideraient le commissariat à mettre sur pied, par exemple, un code, et à s'assurer que ce code est applicable?

Philippe Dufresne: Nous souhaiterions obtenir des pouvoirs d'application de la loi et d'imposition d'amendes. Des sanctions pécuniaires devraient être infligées dans les cas appropriés. Comme je le disais, c'est quelque chose que je ne voudrais pas avoir à utiliser. Je veux que le risque de sanction incite les organisations à investir dans la protection de la vie privée.

Un autre aspect est l'adoption d'un code de pratique. La mise au point d'un code et la mise en place d'un programme entraînent des dépenses pour les organisations. Les codes de pratique constituent un investissement pour une PME, même si elle est relativement grande. Les joueurs de l'industrie et les avocats nous disent entre autres qu'ils ont parfois de la difficulté à convaincre les gestionnaires dans les entreprises parce que ces derniers ne voient pas ce qu'ils retireraient de cet investissement. Ils ne se rallieront pas facilement s'ils n'obtiennent aucune protection réglementaire en échange.

Le projet de loi C-27 énonçait que les organisations qui développent un code et qui le font approuver par le commissaire à la protection de la vie privée pourraient s'en servir pour prouver leur bonne foi lorsque des plaintes sont déposées contre elles.

Le code amoindrirait considérablement les risques d'avoir une amende ou de faire l'objet d'une enquête. À mon avis, cette mesure incitative vise juste, car elle facilite les choses pour les PME et les autres entreprises.

Leslie Church: Merci.

Le président: Merci.

Nous avons largement dépassé le temps alloué, mais je voulais vous donner la chance de répondre à cette question importante.

[Français]

Monsieur Thériault, vous avez la parole pour cinq minutes.

Luc Thériault: Merci, monsieur le président.

Monsieur le commissaire, dans vos remarques préliminaires, vous avez parlé de cette enquête approfondie sur la plateforme de médias sociaux X et son robot conversationnel. Comment cette enquête avance-t-elle? Avez-vous tous les moyens nécessaires pour la faire avancer rapidement?

Compte tenu de la conversation que vous avez eue avec le ministre et des intentions législatives à venir, auriez-vous des moyens accrus pour arriver à un résultat plus rapide?

Philippe Dufresne: Ce que je peux dire, en commençant par la fin de la question, c'est que j'ai clairement dit dans mes discussions publiques et privées qu'il était nécessaire d'avoir plus de pouvoir.

Cela a été démontré très concrètement lorsque j'ai enquêté sur 23andMe, la compagnie de données génétiques. Mon collègue britannique et moi avons fait la même enquête et nous sommes arrivés au même résultat. À la fin, il y a eu une amende dans le monde britannique et, au Canada, il y a seulement eu une recommandation.

Il ressort clairement de cette juxtaposition qu'il est nécessaire que j'aie ces pouvoirs, étant donné qu'il s'agit d'un droit fondamental. Je suis optimiste et je crois que ça viendra, mais j'attends de voir la suite des choses.

Pour ce qui est de notre enquête, nous allons certainement en faire une priorité. Je ne peux pas vous en dire beaucoup en ce qui concerne l'enquête en cours, mais je peux dire ceci: quand nous l'avons entamée, j'ai exprimé ma préoccupation par rapport à ce phénomène et aux risques importants qu'il entraîne pour le droit fondamental à la vie privée des gens, en particulier pour celui des jeunes.

Nous avons les ressources nécessaires et nous allons les utiliser pour donner la priorité à cette enquête.

Luc Thériault: Dans ce processus, à l'heure actuelle, considérez-vous avoir suffisamment les moyens d'enquêter pour arriver à bon port?

Je ne parle pas uniquement de la conclusion de la sanction, mais aussi de la capacité d'affronter ce phénomène qui n'est pas présent que sur ce média social.

• (1720)

Philippe Dufresne: Oui, nous avons nos ressources internes.

Bien sûr, nous pourrions toujours en avoir plus, mais nous sommes conscients des restrictions dans le contexte actuel. Nous travaillons de près avec nos collègues à l'étranger, alors, il y a ce partage d'expertise et de connaissances que nous allons mettre en avant. Nous utilisons les outils que nous avons par rapport aux organisations sur lesquelles nous enquêtons.

Là où nous avons un problème, je le répète, c'est lorsque nous arrivons à la fin de l'enquête, que nous avons constaté une infraction et que l'organisation en cause refuse de mettre en place nos recommandations. Ça a été le cas, par exemple, lors de l'enquête sur Pornhub ou Aylo: nous avons demandé à ce qu'on cesse de publier des vidéos pornographiques sans le consentement de tous les gens qui paraissent dans ces vidéos. L'organisation n'a pas accepté nos recommandations. Je dois donc les amener en cours, ce qui prend du temps et beaucoup de ressources. C'est ça, ma préoccupation, parce qu'en attendant, notre recommandation n'est pas mise en œuvre et les Canadiens ne sont pas protégés dans le cas de ce type de pratique.

Luc Thériault: À propos du phénomène de la diffusion des hypertrucages, à l'heure actuelle, avez-vous des recommandations qui permettraient de mieux encadrer cela?

Philippe Dufresne: J'ai formulé des recommandations, en décembre 2023, avec mes homologues des provinces et des territoires concernant l'intelligence artificielle en général. Cependant, au sujet des hypertrucages en particulier, nous disions qu'il était incorrect de s'en servir pour manipuler des gens, par exemple, pour tenir des propos diffamatoires ou pour diffuser du contenu sexuel sans le consentement de la personne. Il y a donc déjà des balises à cet égard.

Plus récemment, avec le Conseil de la radiodiffusion et des télécommunications canadiennes, ou CRTC, le Bureau de la concurrence et la Commission du droit d'auteur, nous avons publié une analyse des hypertrucages du point de vue de la vie privée, du droit de la concurrence, du droit en matière de radiodiffusion et du droit d'auteur. Dans mon monde, soit celui de la vie privée, les préoccupations sont liées à l'utilisation de cette technologie pour manipuler des gens ou pour diffuser du contenu sexuel sans le consentement de la personne, entre autres. On pense aussi à l'utilisation des données personnelles et à la transparence.

Nous avons donc publié cette analyse en septembre, et c'est un bon résumé. De plus, ça montre la collaboration que nous avons, non seulement avec d'autres commissaires à la protection de la vie privée au Canada et dans le monde, mais aussi avec des organismes de réglementation du contenu numérique, comme le Bureau de la concurrence et le CRTC.

Le président: Merci, messieurs Dufresne et Thériault.

Monsieur Hardy, vous avez la parole pour cinq minutes.

Gabriel Hardy: Merci d'être avec nous, encore une fois.

En répondant à une question de M. Barrett, plus tôt, vous disiez que nous n'aviez pas le pouvoir d'obliger des entreprises comme X à fournir des informations sur la structure de leur fonctionnement ou de les obliger à fonctionner d'une certaine manière.

Quelles recommandations donneriez-vous au gouvernement pour nous permettre d'agir plus rapidement pour protéger la population?

Philippe Dufresne: Je vous remercie de votre question.

En fait, dans le cadre d'une enquête, j'ai le pouvoir de demander de l'information. Alors, pour ce qui est de l'enquête, dans la loi, il n'y a pas de problème.

Là où je n'ai pas de pouvoirs, c'est à la fin de l'enquête. Si je conclus que la loi n'a pas été respectée, je ne peux pas ordonner à l'organisation de cesser une pratique illégale, par exemple.

Gabriel Hardy: C'est ce que vous sembliez dire au sujet de Pornhub, par exemple. Vous lui avez dit ce qu'elle devrait faire, mais c'était une recommandation, donc ce n'était pas réellement contraignant. En fin de compte, vous utilisez beaucoup de moyens, vous faites des recherches et, quand la conclusion arrive, les entreprises font un peu ce qu'elles veulent quand même.

Philippe Dufresne: Dans ce cas-là, c'est ce qui est arrivé.

Heureusement, beaucoup d'entreprises canadiennes avec lesquelles nous faisons affaire acceptent nos recommandations. Il y a un dialogue et elles vont de l'avant. Cependant, quand il y a un refus, c'est problématique. Dans le cas de Pornhub, l'entreprise a refusé d'appliquer deux recommandations qui, à mon avis, étaient essentielles.

Premièrement, nous lui disions de s'assurer d'avoir le consentement direct et clair de tous les gens impliqués. Si c'est indirect, il peut y avoir des cas de pornodivulgateur, par exemple, ce sur quoi portait notre enquête. Deuxièmement, nous voulions aussi que les informations diffusées sur son site de façon inappropriée soient enlevées rapidement et que ce soit facile pour les victimes de les faire enlever.

L'organisation refuse toujours d'appliquer ces deux recommandations, et c'est pourquoi la modification que je demande vise à ce que je puisse délivrer des ordonnances et à ce qu'elles s'appliquent tout de suite. Les gens peuvent contester ma décision devant les tribunaux, c'est normal, ça fait partie de la règle de droit, mais il devrait y avoir une protection immédiate pour les Canadiennes et les Canadiens.

• (1725)

Gabriel Hardy: Ces recommandations ont donc été faites au gouvernement canadien, et vous attendez encore qu'elles soient appliquées.

Avez-vous fait d'autres recommandations dont vous attendez toujours la mise en œuvre? La dernière fois que nous nous sommes vus, vous sembliez dire que votre bureau était un peu sous-financé et que vous n'aviez pas assez de personnel.

Sentez-vous qu'on vous soutient et qu'on vous donne les moyens de vos ambitions? À mon avis, votre objectif étant de protéger les gens, ça ne devrait pas être superflu.

Philippe Dufresne: Nous avons fait plusieurs recommandations de nature substantielle, c'est-à-dire qu'elles requièrent un nouveau projet de loi. D'abord, il y a celle concernant les pouvoirs. Ensuite, il s'agit de reconnaître le droit à la vie privée comme un droit fondamental, de protéger les jeunes et de donner aux gens le droit de faire effacer certaines données qui sont rendues publiques. Il y a aussi la notion d'évaluations préalables et la notion internationale. Quand les données quittent le Canada, on devrait avoir un régime un peu plus rigoureux...

Gabriel Hardy: C'est surtout ça. Vous l'avez dit tantôt et on y revient encore. Combien de recommandations ont été retenues et combien sont en train d'avancer? Alors que le gouvernement Carney est en place depuis un an, combien de vos recommandations sont vraiment en voie d'être mises en œuvre?

Philippe Dufresne: J'attends de voir ce qui sera présenté. Je ne m'attends pas à ce qu'il y ait un projet de loi par recommandation. Je m'attends à ce qu'il y en ait un pour le secteur privé et un pour le secteur public. Pour ce qui est du secteur public, je m'attends à ce que ça progresse par la suite, mais je vais continuer de souligner la nécessité d'agir rapidement à ce sujet...

Gabriel Hardy: À l'heure actuelle, avec vos recommandations, le public n'est pas plus protégé qu'il l'était il y a un an.

Philippe Dufresne: À l'heure actuelle, les recommandations ne sont toujours pas appliquées. Je continue d'utiliser les outils que j'ai du mieux que je peux et en collaborant à l'international, mais il y a ici une occasion qui, à mon avis, devrait être saisie par le Parlement le plus rapidement possible.

Gabriel Hardy: C'est excellent.

Je vais aborder un autre sujet et parler des études que vous avez faites sur l'intelligence artificielle, le vol d'identité, les enjeux de contrôle des données et d'accès à l'information, etc.

On vient d'annoncer un partenariat avec le gouvernement chinois, qui va envoyer 50 000 voitures ici. On sait que les gens passent énormément de temps dans leur voiture. C'est là qu'ils ont des discussions avec leur famille, ils s'en servent pour se déplacer et ils partagent peut-être même des choses. Tantôt, vous avez dit qu'une fois les données sorties du pays, on perd le contrôle.

Voyez-vous là un danger potentiel pour la sécurité des citoyens canadiens?

Philippe Dufresne: Le transfert des données à l'extérieur du Canada est un enjeu qu'il faut contrôler. Je ne crois pas qu'on puisse le freiner complètement, parce que le commerce international est basé là-dessus. Alors, il faut trouver un équilibre entre la protection de la vie privée, la sécurité nationale et la souveraineté numérique, tout en permettant le commerce international. À cet égard, les mesures en matière de protection de la vie privée sont utiles, parce que ça amène des éléments liés au consentement et à la transparence.

Sait-on ce qui se passe? Accepte-t-on que ça se passe? Est-il facile d'opposer des refus? Les fins sont-elles appropriées? Il est important de vérifier tout ça dans le contexte des voitures numériques, tant au Canada qu'à l'étranger, mais encore plus lorsque les données quittent notre territoire.

Le président: Merci, monsieur Dufresne.

Madame Lapointe et monsieur Hardy, vous avez l'occasion de poser d'autres questions, si vous le voulez.

Madame Lapointe, vous avez la parole.

Linda Lapointe: Merci beaucoup, monsieur le président.

Je souhaite la bienvenue aux deux témoins.

C'est très intéressant. En ce qui concerne le travail du gouvernement sur vos recommandations, je crois que vous avez bon espoir de voir vos recommandations aboutir, n'est-ce pas?

Philippe Dufresne: J'ai bon espoir. Le ministre de l'Intelligence artificielle a fait des déclarations publiques dans lesquelles il a dit qu'on y travaillait et que l'intelligence artificielle en avait besoin. Ça revient à ce que je dis: il faut s'assurer d'avoir cette innovation et une économie forte, tout en protégeant la vie privée. C'est possible de le faire.

Linda Lapointe: Merci.

Tantôt, vous avez parlé de votre homologue de la Grande-Bretagne. Vous disiez que vous aviez fait la même enquête sur une entreprise de données génétiques, c'est-à-dire une entreprise qui reçoit par la poste des tests de gens qui veulent savoir qui sont leurs ancêtres, et que vous étiez arrivés au même résultat, mais qu'en Grande-Bretagne, des sanctions avaient pu être infligées à cette entreprise.

Philippe Dufresne: C'est exactement ça. Il s'agissait d'une entreprise mondiale de généalogie génétique, et il y a eu une atteinte à la vie privée. Des gens ont réussi à voler les données de millions de personnes, dont 300 000 Canadiens. Imaginez ce que ça représente. Les données génétiques sont extrêmement sensibles.

Quand nous recevons des plaintes de cette nature — dans ce cas-ci, j'ai mené l'enquête avec mon collègue britannique —, nous vérifions si l'entreprise a traité les données de façon suffisamment prudente et s'il y avait suffisamment de mécanismes de protection. Malheureusement, nous avons constaté que les mots de passe et les systèmes de protection n'étaient pas assez robustes et qu'on avait été trop lent à réagir aux indices montrant que des malfaiteurs tentaient d'entrer.

Mon collègue britannique et moi sommes donc arrivés à la même conclusion. Cependant, lui, il a le pouvoir de délivrer des ordonnances et de donner des amendes quand il y a de telles violations importantes, ce qu'il a fait. Je pense qu'il s'agissait de millions de livres. Quant à moi, je ne pouvais que faire des recommandations. La situation était donc assez flagrante. Nous avons fait une conférence de presse et on m'a demandé pourquoi je n'avais pas imposé d'amendes. J'ai répondu que je n'avais pas le pouvoir de le faire.

• (1730)

Linda Lapointe: Étant donné qu'il y a eu des sanctions en Grande-Bretagne, des correctifs ont-ils été apportés plus rapidement par cette entreprise?

Philippe Dufresne: En fait, l'entreprise a accepté toutes nos recommandations. Ça montre que, quand on a le pouvoir de donner des amendes et de délivrer des ordonnances, ça encourage l'organisation à accepter les recommandations. Est-ce qu'elle l'aurait fait s'il n'y avait eu que moi? C'est possible. Comme je vous l'ai dit, nous pouvons être très persuasifs. Nous critiquons publiquement les organisations, et elles n'aiment pas ça. Cependant, s'il n'y a aucune conséquence financière, c'est difficile de convaincre les conseils d'administration de dépenser des sommes élevées pour appliquer les recommandations. Selon moi, ça prend un incitatif. Je me dis que tout dirigeant d'entreprise va fixer son attention sur les risques, qu'ils soient juridiques ou financiers.

Linda Lapointe: Plus tôt, vous avez beaucoup parlé de l'intérêt supérieur de l'enfant et du consentement. Je sais que, maintenant, il y a une tendance à ne pas afficher des photos des enfants parce qu'on n'a pas leur consentement. En effet, lorsque les enfants sont petits, ils ne peuvent pas acquiescer.

Comment suggérez-vous qu'on s'assure du consentement et de l'intérêt supérieur des enfants en ce qui a trait à toutes les photos qui sont affichées sur les médias sociaux?

Aussi, comment faut-il tenir compte de la reconnaissance faciale et des possibilités qu'elle prenne une mauvaise tangente pour ces enfants, qui n'ont pas choisi d'être sur les médias sociaux?

Philippe Dufresne: Tous ces exemples expliquent pourquoi j'ai fait de la protection de la vie des enfants l'une de mes trois priorités stratégiques. En fait, c'est la raison pour laquelle ils sont vulnérables: parfois, on va mettre leurs données en ligne, sans qu'ils le sachent, sans leur consentement. Alors, il faut faire plusieurs choses. Premièrement, il faut interpréter la loi de façon à tenir compte de leur intérêt supérieur. On le fait en droit de la famille. On le fait dans tous les domaines du droit, en général, mais, dans le contexte de la vie privée, on ne le fait pas tout à fait suffisamment.

Par exemple, il s'agirait de dire que le consentement d'un enfant doit être éclairé et approprié à leur âge. Dans certains cas, ce sont les parents qui le donneront et, dans d'autres, ils pourront le faire eux-mêmes. Il faut que la communication soit adaptée à l'âge de l'enfant. Je viens de créer un conseil consultatif des jeunes pour mon commissariat. Je vais rencontrer des jeunes de 13 à 18 ans. Je vais leur demander comment ils se sentent, ce qu'ils veulent en matière de protection de la vie privée et quelle utilisation des médias sociaux ils font.

En résumé, il faut en faire davantage.

Linda Lapointe: J'aimerais poser une autre question.

Le président: Votre temps de parole est écoulé. Vous pourrez poser d'autres questions lors d'un autre tour.

Nous allons maintenant nous tourner vers M. Hardy, qui va partager son temps de parole avec M. Cooper.

Monsieur Hardy, vous disposez de cinq minutes.

Gabriel Hardy: Ma question est assez simple, et elle revient à celle que je vous ai posée plus tôt sur le gouvernement chinois.

En fait, jusqu'à très récemment, la Chine était identifiée comme un des plus grands dangers pour notre nation. Par contre, maintenant, on entame le commerce avec eux. Vous travaillez à l'échelle internationale, avec des collègues d'un peu partout au monde, comme en Australie, en Angleterre, etc.

Vous sentez-vous à l'aise de travailler avec la Chine si une collaboration s'entame? Avez-vous de bons contacts? Sentez-vous que c'est une nation avec laquelle on peut avoir des échanges tout en estimant que la sécurité des données personnelles sera aussi bien assurée que par vos collègues d'Angleterre, d'Australie et peut-être même des États-Unis?

Philippe Dufresne: Le Canada est très actif sur la scène internationale et il y a beaucoup de réseaux. Donc, il y a le Global Privacy Enforcement Network, la Table ronde des autorités de protection des données et de la vie privée du G7, dont nous avons parlé, et il y a le Forum des autorités de protection de la vie privée de la zone Asie-Pacifique. Nous travaillons donc de près avec le Japon, la Corée et Singapour. Il y a des commissaires aussi à Hong Kong avec

qui nous avons des échanges. Je dois dire que ce sont de bons échanges en matière de principes de vie privée, sur ce qu'on veut voir en tant que commissaires à la protection de la vie privée.

Je n'ai pas eu d'échanges avec le gouvernement chinois, par contre, mais simplement avec les vis-à-vis en ce qui a trait à la protection de la vie privée. On essaie d'avoir une voix commune, la plus forte possible, ce qui nous a permis de faire plusieurs déclarations à cet égard. Cela étant dit, dans notre enquête sur TikTok, par exemple, nous avons émis des préoccupations, à savoir que le gouvernement chinois a des pouvoirs accrus pour obtenir de l'information de compagnies privées, et nous voulions que TikTok signale qu'il y avait une possibilité que les données parviennent en Chine et soient accessibles dans ce contexte.

• (1735)

Gabriel Hardy: La Chine a donc travaillé avec vous.

Philippe Dufresne: Oui, TikTok a accepté de modifier ses politiques là-dessus.

Le président: Il vous reste moins de trois minutes, monsieur Cooper.

[Traduction]

Michael Cooper: J'aimerais revenir au projet de loi du gouvernement sur la cybersécurité. En plus de dispenser le ministre de l'exigence d'obtenir une autorisation judiciaire pour ordonner à un fournisseur de services de télécommunications de divulguer des métadonnées, des renseignements sur les comptes des abonnés, des listes de sites Web consultés, l'emplacement de données financières et d'autres données personnelles, le projet de loi, comme vous l'avez mentionné, et les pouvoirs qu'il octroie permettent de faire l'impasse sur les critères de nécessité et de proportionnalité. La pertinence est la norme en vigueur actuellement. Est-ce exact?

Philippe Dufresne: C'est le cas pour un des pouvoirs. Le projet de loi en renferme plusieurs. Certains...

Michael Cooper: Je parle des pouvoirs octroyés par l'article 15.

Philippe Dufresne: Oui.

Michael Cooper: C'est une norme très permissive.

Philippe Dufresne: Cette norme est beaucoup plus permissive que le critère de nécessité et de proportionnalité, d'où ma recommandation de remplacer l'une par l'autre.

Michael Cooper: Aucune autorisation judiciaire et aucune norme de nécessité et de proportionnalité ne sont en place. Pourtant, ce sont des données et des renseignements très importants touchant à la vie privée des Canadiens qui sont en cause.

Ce sont à mon avis des prétextes utilisés par le gouvernement pour faire taire les préoccupations liées à la protection des renseignements personnels suscitées par le projet de loi. Le gouvernement affirme que les attentes en matière de protection de la vie privée sont faibles lors de la conduite d'activités réglementées. C'est vrai. Les attentes en la matière sont moindres ou raisonnables.

Je voudrais confirmer que les activités liées à Internet ne sont pas des activités réglementées. En fait, ces activités devraient susciter des attentes élevées en matière de protection de la vie privée.

Philippe Dufresne: Dans le cas d'Internet, je n'ai pas réclamé d'exigence relative à l'obtention d'un mandat ou d'une autorisation judiciaire. J'ai réclamé, et je réclame toujours un critère plus strict, qui devrait être, à mon avis, la nécessité et la proportionnalité.

Michael Cooper: Je suis d'accord, mais je voulais également faire valoir que l'utilisation d'Internet, des métadonnées et de toutes ces choses est une activité, selon les jugements de la Cour suprême, qui suscite des attentes élevées en matière de protection de la vie privée. Cette affirmation est-elle juste?

Philippe Dufresne: J'estime que la mise en place d'une norme de nécessité et de proportionnalité, qui oblige à déterminer quelles données sont nécessaires et en quelle quantité, serait une mesure suffisante. Il faut dire aux organisations de ne pas recueillir plus de renseignements que ce dont elles ont besoin, et les encourager à mettre en place des mécanismes ayant une incidence minimale sur la vie privée.

Le président: Merci, monsieur Cooper.

[Français]

Madame Lapointe, je crois que vous allez partager votre temps de parole.

Linda Lapointe: Oui, monsieur le président, mais j'interviendrai brièvement avant de le faire.

Le projet de loi C-16, Loi modifiant certaines lois en matière pénale et correctionnelle (protection de l'enfance, violence fondée sur le sexe, délais et autres mesures) a été déposé à la Chambre. Cela englobe aussi les images hypertruquées ainsi que le partage non consensuel d'images intimes.

Êtes-vous au courant de ça?

Philippe Dufresne: Oui, je suis au courant.

Il s'agit là du domaine du droit criminel. Pour moi, la Loi sur la protection des renseignements personnels continue de jouer un rôle important. Ce ne sont pas tous les enjeux qui vont nécessairement pouvoir être traités par le droit criminel, alors cette protection est importante.

Comme je l'ai fait dans le cadre de mon enquête sur Aylo, je pense qu'on peut aussi utiliser des lois non criminelles pour protéger les Canadiennes et les Canadiens.

Linda Lapointe: Nous avons reçu plusieurs témoins très alarmistes, qui disaient que la fin du monde allait arriver à cause de la superintelligence, de l'intelligence artificielle. Certains comparent ça à une guerre nucléaire ou au changement climatique.

Vous avez rencontré les autres représentants du G7 qui sont des commissaires comme vous. Les gens, ailleurs dans le monde, sont-ils alarmistes eux aussi?

• (1740)

Philippe Dufresne: Lors des événements internationaux auxquels je participe, y compris au G7, il y a toutes sortes de discussions. Sur le plan de la protection de la vie privée, mes collègues et moi nous concentrons sur la situation qui est devant nous en ce moment et sur l'avenir, comme les agents de l'intelligence artificielle qui commencent à être parmi nous.

En effet, nous entendons ces préoccupations. Lorsque nous étions au G7, nous avons invité M. Yoshua Bengio à venir nous parler de certaines préoccupations. Je pense qu'il est utile d'avoir ces échanges. Cependant, ce n'est pas unanime. Autrement dit, tout le monde n'est pas pessimiste. Le message que nous entendons porte sur ce que nous devons faire pour nous protéger maintenant. Renforcer la protection de la vie privée, avoir des éléments comme la transparence, le contrôle humain, le consentement et les fins légi-

times, e tout ça rendront, à mon sens, beaucoup moins probable cette évolution que nous ne voulons pas voir.

Le président: Il vous reste deux minutes quarante-cinq secondes, monsieur Sari.

Abdelhaq Sari: Merci beaucoup, monsieur le président.

De manière générale, le droit au savoir ou le droit à la cognition est fondamental, et il me tient vraiment à cœur. Il s'agit de l'un des risques auxquels je m'attarde dernièrement. J'ai évoqué, tout à l'heure, la question de la concentration du pouvoir cognitif, comme on l'appelle. La concentration du pouvoir cognitif, malheureusement, n'est pas uniquement sur le savoir lui-même, mais sur l'élaboration et la création de ce qu'on appelle la connaissance créatrice. Il ne s'agit pas uniquement de la connaissance elle-même, mais de la création de la connaissance.

Ce qui m'intrigue et m'inquiète aussi n'est pas uniquement la croissance accrue de ce pouvoir de concentration, mais aussi le fait qu'elle est concentrée entre les mains de quelques acteurs et non pas de plusieurs. Quand on dit le mot « universel », ça ne veut pas nécessairement dire que c'est international et que c'est le cas dans tous les pays. La croissance accrue du pouvoir est concentrée entre les mains de quelques acteurs.

Vous êtes le commissaire et nous sommes les législateurs. Quel cadre juridique avons-nous actuellement?

Quelle trajectoire devons-nous suivre afin d'être au moins conscients du problème et de ce que nous devrions faire pour protéger non seulement les renseignements personnels, mais également ce droit au savoir de la meilleure façon, et afin qu'il ne soit pas dirigé d'une seule façon?

Philippe Dufresne: Je pense qu'il nous faut utiliser tous les outils que nous avons. Dans mon cas, cela consiste à travailler avec plusieurs autres disciplines, donc dans mon domaine de protection de la vie privée. La psychologie cognitive devient aussi très importante, c'est-à-dire le fait de comprendre comment on prend des décisions et comment on ne prend pas des décisions, comment on peut être manipulé par des algorithmes, par des messages, par l'intelligence artificielle. Il s'agit de s'assurer du rôle que joue l'humain dans tout ça.

J'ai fait des recommandations sur le plan législatif, à savoir qu'on veut une plus grande transparence, qu'on veut la protection de la vie privée, donc de la dignité humaine dès le départ, dès la conception.

Pour ce qui est du petit nombre de grands acteurs, je pense que ça vient souligner l'importance de la collaboration internationale aux échelons gouvernemental, parlementaire et réglementaire, et nous le faisons. Un seul pays aurait de la difficulté à réglementer ça, parce que ça traverse les frontières. Cependant, en agissant de concert avec la communauté internationale, nous pouvons y arriver, et c'est ce que nous faisons.

Le président: Merci, monsieur Sari.

[Traduction]

Voilà qui met fin aux témoignages.

Monsieur Dufresne, au nom du Comité, je vous remercie d'avoir été des nôtres. J'ai confiance en votre capacité à gérer à la fois ce dossier et votre personnel. Vos témoignages devant le Comité raffermissent encore plus cette confiance. Merci beaucoup des réponses que vous avez fournies aujourd'hui.

Il y aurait quelques points dont je voudrais discuter avec le Comité, notamment deux questions urgentes à régler.

Tout d'abord, j'ai besoin de l'assentiment de la Chambre au sujet de la Loi sur les conflits d'intérêts. Le projet de rapport devrait être déposé sur nos bureaux bientôt. Il me faut une motion à cet effet. Nous n'avons pas encore obtenu l'approbation qui nous permettra de présenter le rapport à la Chambre.

Ensuite, il me faut un consentement concernant la Loi sur le lobbying et l'étude à venir sur le sujet. Je vais en parler jeudi. Il me faudra pour chacune des deux questions une motion du Comité à soumettre à la Chambre pour nous assurer que le Comité est la tribune appropriée et que la Chambre est d'accord pour que nous examinions les deux questions. Ce sont les points dont je voulais vous faire part.

Le ministre de l'IA viendra témoigner jeudi. Une motion a été proposée pour qu'il témoigne pendant deux heures, mais sa visite ne durera qu'une heure. Des efforts soutenus ont été faits pour que le ministre soit présent pendant les deux heures, mais son témoignage ne durera qu'une heure, suivi du témoignage des fonctionnaires pendant la deuxième heure. Je voulais que vous sachiez que la greffière a déployé tous les efforts possibles pour que le ministre témoigne deux heures conformément à la demande formulée dans une motion du Comité adoptée à l'unanimité.

C'est tout. Je n'ai pas d'autres questions. Je vous souhaite une bonne journée.

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>