



< >

Déclaration de Montréal IA responsable_

</ >

PARTIE 5

RAPPORT DE LA COCONSTRUCTION EN LIGNE ET DES MÉMOIRES REÇUS



Ce document est une partie du
**RAPPORT DE LA DÉCLARATION DE MONTRÉAL
POUR UN DÉVELOPPEMENT RESPONSABLE
DE L'INTELLIGENCE ARTIFICIELLE 2018.**
Vous retrouverez le rapport complet [ICI](#).

RÉDACTION

MARTIN GIBERT, conseiller en éthique pour
IVADO et chercheur au Centre de recherche
en éthique

Dans ce document, l'utilisation du
genre masculin a été adoptée afin de
faciliter la lecture et n'a aucune intention
discriminatoire.

TABLE DES MATIÈRES

1. INTRODUCTION	220
2. LE QUESTIONNAIRE EN LIGNE	221
Bien-être (environnement, prudence)	221
Autonomie	225
Justice (équité, solidarité, diversité)	229
Vie privée (intimité)	234
Connaissance (publicité, prudence)	238
Démocratie (publicité, diversité)	241
Responsabilité (prudence)	244
3. SYNTHÈSE DES MÉMOIRES REÇUS	248
Vie privée	249
Justice	251
Responsabilité	252
Bien-être	253
Autonomie	255
Connaissance	256
Démocratie	256
CRÉDITS	I
PARTENAIRES	II

1. INTRODUCTION

En novembre 2017, au terme d'un congrès organisé par l'Université de Montréal au Palais des congrès de Montréal, était lancée la première étape de la *Déclaration de Montréal pour un développement responsable de l'IA*. Une version préliminaire de cette Déclaration articulée autour de 7 principes allait servir de base à une phase de coconstruction débouchant sur une nouvelle version. Si les ateliers de discussion ont permis de consulter les citoyens et experts, d'autres moyens de participer à la réflexion collective étaient possibles : 1) en répondant à un questionnaire en ligne accessible sur le site de la Déclaration (www.declarationmontreal-iaresponsable.com), et 2) en envoyant des mémoires sur un ou plusieurs aspects de la Déclaration. Le présent rapport propose la synthèse des mémoires reçus et des réponses au questionnaire. Le rapport sur les ateliers de coconstruction est également disponible sur le site web de la Déclaration.

2. LE QUESTIONNAIRE EN LIGNE

Le questionnaire en ligne était composé de 35 questions, cinq pour chaque principe. Il a retenu l'attention de 83 répondants, dont 17 anglophones. Comme on pourra le voir dans la synthèse, plusieurs ont une connaissance avancée de l'IA et des enjeux éthiques et sociaux que soulève son développement.

Les questions sont présentées dans l'ordre du questionnaire qui reprend le plan de la Déclaration préliminaire. La Déclaration révisée étant plus complète (elle se compose de dix principes), les nouveaux principes pertinents ont été ajoutés entre parenthèses à ceux de la version préliminaire.

BIEN-ÊTRE (ENVIRONNEMENT, PRUDENCE)

1. COMMENT L'IA PEUT-ELLE CONTRIBUER AU BIEN-ÊTRE ?

Une question générale qui suscite beaucoup de réponses, et des réponses très variées. On invoque beaucoup les espoirs pour la santé et l'aide aux personnes âgées ou en situation de handicap. L'IA semble aussi un espoir pour diminuer les impacts environnementaux même si on note que « le développement de l'IA a une empreinte écologique

(donc un coût direct sur le bien-être) souvent négligée, bien que non négligeable ». Plusieurs notent que l'IA pourrait remplacer l'humain dans des tâches à risque. L'aspect « aide à la prise de décision » est aussi mentionné plusieurs fois, en particulier sous la forme d'un possible assistant personnel qui pourrait aussi nous assister dans nos recherches d'information.

On s'attend aussi à ce que l'IA améliore la productivité et nous libère des tâches répétitives et routinières. Elle pourrait aussi anticiper nos attentes et nos besoins ou simplement passer l'aspirateur à notre place. Une disposition importante : l'IA améliorera notre bien-être « à condition que nous vivions dans une véritable démocratie, qu'elle soit au service de tous et pas seulement au service de quelques privilégiés ».

EXTRAITS CHOISIS

"The AI or any technology will create a lot more value for the rural population than the urban population. A single smartphone can provide immense value, and anything that can collect data is a breeding ground for AI: Better education, better farming technology (Eg: crop analysis, robot-farming)."

2. EST-IL ACCEPTABLE QU'UNE ARME AUTONOME PUISSE TUER UN ÊTRE HUMAIN ? UN ANIMAL ?

Une très forte majorité des gens répondent non à cette question, souvent avec beaucoup d'insistance et de points d'exclamation. On invoque que « tuer doit rester dans les mains de l'humain qui doit avoir pleine conscience de son geste ». L'idée d'interdire légalement les systèmes d'armes autonomes est d'ailleurs mentionnée plusieurs fois. Un répondant souligne aussi le risque d'une course à l'armement et la possibilité d'erreurs de programmation. Quelques répondants font une distinction : non pour un humain, oui pour un animal (« dans un contexte de régulation de population »). Certains cas semblent

faire exception : une machine tuant un condamné à mort ou un « tigre qui s'échappe de sa cage et met en danger le public ». Dans tous les cas, il apparaît que l'IA ne devrait être qu'un outil de mise à mort dont la responsabilité incombe à l'humain. Un répondant développe toutefois une réflexion critique et soulève une question pertinente : « si cette arme peut prendre une meilleure décision qu'un humain, pourquoi pas ? ».

On remarque aussi que cette question dépend beaucoup du contexte : "It would be acceptable for an autonomous weapon to kill a human being or an animal in any circumstances where it would be acceptable for a human or other creature to kill a human or animal."

EXTRAITS CHOISIS

« Les armes autonomes ne devraient pas exister, elles devraient être bannies au même titre que les armes chimiques. Un humain devrait toujours être aux contrôles d'une arme : il aura ainsi la responsabilité morale de son geste. »

"No! (...) a horrific scenario could ensue from an unethical manufacturer or rogue programmer who perhaps, unbeknownst to either the weapon's company or weapon purchaser, may secretly design, code & program autonomous weapons which reflect their secret views & biases as a Neo-Nazi or KKK supporters, for example."

"Why would you think a HUMAN should be able to kill somebody? If you have a reason, then why doesn't it apply to an AI. There's no reason humans should always occupy a

privileged position with respect to killing other humans. Obviously "AI" at the moment is not even ready for consideration for this, but that's unlikely to be permanent (assuming you think anything or anybody should be killing whoever or whatever you're thinking about killing). The interesting question may be how you'll know when it's changed, and how you manage the transition."

3. EST-IL ACCEPTABLE QU'UNE IA CONTRÔLE UN ABATTOIR ?

Comme à la question précédente, une large majorité répondent que non (on note toutefois plus de réponses favorables). On retrouve l'argument de la banalisation de la violence par la distance psychologique invoquée dans la question précédente : « Cela éloignerait encore plus l'homme de cette action de tuer l'animal », ou encore : « Il ne faut pas offrir une nouvelle couche de couardise à l'humain derrière laquelle il peut se cacher en déléguant une tâche moralement répréhensible à un robot. » L'argument environnemental est aussi invoqué : « Ce n'est pas l'orientation à privilégier pour l'avenir de la planète et des humains qui l'habitent. Je préférerais que l'IA contrôle des serres et des édifices zéro émission de CO2 et de déchet. »

Quelques arguments en faveur d'un tel projet : éviter la maltraitance, le stress des animaux et améliorer l'hygiène. (Un répondant explique néanmoins qu'un humain devrait toujours superviser les opérations d'abattage précisément pour éviter la cruauté...). Une position mixte : l'IA pourrait contrôler le découpage et l'emballage, mais pas l'abattage. Plusieurs se demandent toutefois si les abattoirs sont acceptables, même sans IA.

« Je comprends que cela soulagerait les responsables de ces tâches morbides cependant cela ne serait absolument pas éthique. »

« Les conditions d'abattage seraient peut-être légèrement améliorées, mais la pratique elle-même durerait plus longtemps, puisque l'on pourrait plus facilement s'en détourner. »

"Yes, to the extent that the AI follows humane (and human) protocol."

"Interesting question. One side of me says yes, the other no. What we are doing to other animals that we raise for food already has some serious ethical issues. When I read about the life of the average chicken raised for food, I was shocked. Totally automating the process of raising food, including having AI do the killing would just put the fate of these animals even more out of sight, out of mind. So, on balance, I think I am against an AI controlled abattoir."

"Slaughterhouses already exist and won't stop existing anytime soon. AI can make sure that the method of slaughter is ethical and is done in the most humane way possible. This can also strictly ensure and maintain safety standards."

4. DEVRAIT-ON CONFIER LA GESTION D'UN LAC, D'UNE FORÊT OU DE L'ATMOSPHÈRE TERRESTRE À UNE IA ?

Cette question suscite beaucoup de méfiance à l'égard de l'IA, mais aussi des espoirs de solutions à la crise environnementale. Il ressort de cela encore une fois que l'IA qui prendrait soin de l'environnement devrait être « configurée par des êtres humains responsables qui ont à cœur la préservation ». Certains y voient même un espoir, en particulier pour le climat, mais dans une logique de coopération humain/IA plutôt que d'une délégation du problème à l'IA. Plusieurs répondants espèrent une IA qui ne serait pas corruptible et dans une logique de recherche de profit.

On note le risque de détournement malveillant d'une IA qui aurait une telle mission. Et une pointe de cynisme : « l'humain détruit tous les milieux naturels auxquels il touche ou presque, alors ça ne pourrait pas être pire... » Surgit aussi le thème du remplacement de l'humain : « On pourrait tout laisser faire à l'IA, mais reste à savoir si l'on souhaite que l'homme devienne assisté sur tout son potentiel. » Et un principe démocratique : aucune entité humaine ou artificielle ne devrait pouvoir décider seule de la gestion de l'environnement — cela devrait venir d'une coopération de tous les humains.

EXTRAITS CHOISIS

« Non, car nos connaissances sont largement insuffisantes pour juger des répercussions à long terme des actions décidées par l'IA. »

"It depends on what the instructions given to the IA and how much absolute control it holds. I think an AI trained on environmental systems and with ability to monitor and consider big (environmental) data could make much better decision than any group of individuals, effectively helping to protect the environment and regenerate those

that may have been affected by industry, etc.”

« Cela peut se faire avec l’assistance d’une IA ; mais l’IA ne sait pas d’elle-même ce qui est bon pour le lac, la forêt ou l’atmosphère : elle est plus performante du point de vue de la rationalité instrumentale, mais ne peut se donner à elle-même ses propres fins. »

“At present, AI would be most useful in the collection and analysis of data.”

“Eventually, machines may be more competent than people to make almost all decisions. But, if we give the machine control and stop monitoring what and how it does what it does, the ability of human beings to manage our affairs will pass out of the living memory of humans, and we will be entirely dependent upon machines. This does not seem to me to be a good future for human beings.”

DEVRAIT-ON DÉVELOPPER DES IA CAPABLES DE RESSENTIR DU BIEN-ÊTRE ?

Beaucoup d’hésitations et d’intuitions contradictoires ici. Est-ce seulement possible, techniquement parlant ? On mentionne que la sentience pourrait permettre à l’humain de contrôler ou de punir des IA. D’autres y voient un intérêt afin qu’une IA puisse mieux comprendre les humains (et autres êtres sentients) et faire de l’empathie avec

eux. L’intelligence émotionnelle semble aussi requise pour faire de bons jugements moraux. Toutefois, une empathie simulée paraît suffisante : car on voit un danger à ce qu’une IA sentiente privilégie son bien-être personnel aux fonctions que l’humain lui aurait assignées. « L’IA doit rester un instrument au service de l’humain, pas un simili-humain ». Et plusieurs se demandent « à quoi bon ? » Et un répondant s’inquiète pour l’IA : « Je préfère que l’IA comprenne le bien-être plutôt qu’elle le ressente, surtout que dans la notion de bien-être il y a aussi la notion de mal-être. »

EXTRAITS CHOISIS

“I think it makes most sense to approach the development of general AI as the development of a calculator/tool. Developing a personified, sentient AI may be bringing new life to the world. I’m sure it would be treated fairly or with rights.”

« Question des plus complexes. Si elle ressent du bien-être, cela va générer le désir de le maximiser. Utile dans une logique de récompense des apprentissages. Toutefois, va-t-elle balancer son bien-être de machine avec celui des humains et des êtres vivants ? »

« Oui, mais seulement en fonction de l’accomplissement des tâches qui lui sont dévolues. Il serait ainsi possible de développer des IA qui, grâce à cette satisfaction du devoir accompli, chercheraient à s’améliorer constamment, mais seulement dans le domaine spécifique où elles opèrent. »

AUTONOMIE

1. COMMENT L'IA PEUT-ELLE CONTRIBUER À L'AUTONOMIE DES ÊTRES HUMAINS ?

L'IA présente une relation ambiguë avec l'autonomie : elle nous rend dépendants d'elle-même (« on ne pourra plus se dissocier de l'IA »), tout en libérant l'humain de certaines tâches cognitives aliénantes (ex. : conduire une voiture, démarches administratives), voire de la nécessité de travailler. Dans les commentaires, c'est toutefois l'aspect positif et libérateur qui ressort le plus. Le modèle du partenariat est une option, tout comme celui de l'IA simple assistante. Et de grands espoirs sont possibles : l'IA pourrait améliorer la condition humaine et tout particulièrement les personnes en situation de handicap. On note aussi l'espoir d'une médecine moins invasive et qui vienne en aide aux personnes âgées en perte d'autonomie.

EXTRAITS CHOISIS

« L'IA devrait être utilisée pour redonner de l'autonomie (physique ou mentale) à des personnes handicapées. Seule une personne apte à contrôler entièrement la configuration des algorithmes d'un système d'IA pourrait gagner de l'autonomie, tous les autres en perdent parce qu'ils se fient à des décisions qui sont prises par quelqu'un d'autre. »

« Aucun système n'aidera (et n'aide aujourd'hui) l'autonomie des humains s'il vient de l'entreprise privée. La réglementation et l'implication du secteur public sont essentielles au maintien de l'équilibre. »

« En les libérant des tâches qu'ils ne souhaitent pas faire, et en améliorant leur état émotionnel et leur compréhension du monde. »

"HUMANS will deliberately develop AI to force other humans to follow their values or act according to their interests. And those humans will see themselves as benevolent in doing that which is the really scary part."

"AI can automate most of the trivial things that we spend a lot of time doing. Almost everything that we do without actively thinking about it can in a way be simplified or made more convenient using AI. But this also has to ensure that humans don't become too dependent on the technology, which would then handicap their life instead of providing more autonomy."

2. FAUT-IL LUTTER CONTRE LE PHÉNOMÈNE DE CAPTURE DE L'ATTENTION DONT S'ACCOMPAGNENT LES AVANCÉES DE L'IA ?

Le phénomène suscite beaucoup de scepticisme (« j'ai besoin de plus d'information »). Mais plusieurs notent le risque d'une « hypnose technologique », « surtout chez les adolescents ». Un répondant résume : « il ne faut pas devenir les esclaves de nos technologies ». Et un autre propose de soigner l'IA avec de l'IA : « Il faudrait savoir dans quels buts est captée l'attention. Les buts mercantiles, qui sont le moteur d'applications comme Facebook devraient pouvoir être freinés par d'autres applications IA, sortes de contre-mesures, mises à la disposition des utilisateurs pour lutter contre les intrusions. »

Évidemment, capter l'attention des gens sur des problèmes éthiques paraît une bonne idée : « c'est

un moyen comme un autre de s'assurer que ces débats auront lieu. » Et une remarque pleine de bon sens : "Yes, businesses should be prevented from manipulating people's attention in ways that those people don't control or understand. That's not intrinsically related to AI; it's just that AI is a convenient, powerful, and therefore dangerous tool for it."

EXTRAITS CHOISIS

« Oui. En éduquant les populations dans un premier temps. Puis en légiférant pour imposer un cadre opérationnel s'inscrivant dans des valeurs humanistes (vérité, justice, bonté, respect, etc.) »

"All technologies, from radio frequencies, nuclear energy to cryptography must live within a regulatory framework. Attention seeking AI could be classified as addictive entertainment, like gambling."

"One way to combat this is awareness about the problem, the fact that this is happening is not known to many (hypothesis). And give the user proper tools to combat this: nudge the user to actually learn the skill using small dopamine hits until the user doesn't need it anymore."

3. FAUT-IL S'INQUIÉTER DE CE QUE DES HUMAINS PRÉFÈRENT LA COMPAGNIE DES IA À CELLE D'AUTRES HUMAINS OU D'ANIMAUX ?

Aucune tendance ne se démarque de façon claire. Certes, on s'inquiète de ce que la technologie sépare ou isole les humains : « l'humain doit rester un être social » et il faut prendre garde à ce que l'humain n'oublie pas ses compétences sociales comme l'empathie. Mais de ce point de vue, l'IA « ne serait pas pire que les jeux électroniques ». Des études de psychologie seront sans doute nécessaires pour évaluer les risques d'une nouvelle dépendance. Mais on y voit aussi un bénéfice potentiel pour les personnes seules ou pour certains profils psychologiques : des enfants autistes, par exemple, peuvent avoir plus de facilité à communiquer avec une IA qu'avec un humain. La chose devrait toutefois rester marginale, car « si un humain ne veut plus de contact avec d'autres humains, l'humanité disparaît ». Mais attention au paternalisme : si cela ne nuit pas aux autres, pourquoi empêcher des relations fortes entre humain et IA. Et on accepte en général que certaines personnes préfèrent la compagnie des animaux à celle des humains.

EXTRAITS CHOISIS

« Si les agents IA n'ont pas de sensibilité ni de sentiments, ils ne sont pas de bonne compagnie. Moins que des animaux même. Ce sont des choses, pour l'instant. Des machines. »

« Si la personne n'a pas d'autre option, ça peut être une bonne chose. Sinon on va commencer à avoir des difficultés à vivre en communauté. »

« Non, beaucoup d'êtres humains sont déjà captifs de relations avec des objets ou encore avec des personnages fictifs (télévision,

téléromans, “amis” des réseaux sociaux). La compagnie d’une IA aurait au moins l’avantage de présenter un certain degré d’interactivité qui pourrait s’avérer particulièrement bénéfique chez les personnes âgées ou seules. »

“This is a legitimate concern. It can be compared to the preference for texting as a substitute for direct human-to-human interaction.”

“If you care about people’s autonomy, then LET THEM MAKE THEIR OWN DECISIONS. It doesn’t matter whether you’re “worried”, because it’s purely none of your business, full stop.”

“Even if technologies like VR [virtual reality] are developed to an almost realistic level, it would only increase social isolation, and it would be detrimental in the long run. Social security, i.e the fact that there are people to support you and will be there with you in your time of need, is invaluable!”

4. PEUT-ON DONNER SON CONSENTEMENT ÉCLAIRÉ FACE À DES TECHNOLOGIES AUTONOMES DE PLUS EN PLUS COMPLEXES ?

Cela va être difficile pour deux types de raison, soulignent plusieurs répondants : la complexité des machines et la complexité des clauses légales. Personne ne lit les termes de consentement des applications ou des plateformes qui sont trop complexes (jargon juridique) : lorsque les gens

« acceptent » ont-ils réellement le choix ? Ces consentements ne peuvent pas être considérés véritablement éclairés. “How often do we sign off on online agreements saying we read them when we didn’t?”

Il reste donc à créer des systèmes qui donnent confiance et soient sécuritaires. Le manque de littératie numérique est aussi pointé, ainsi que la nécessité d’y remédier par l’éducation. Cela montre d’ailleurs « l’importance d’établir un code d’éthique sur lequel l’IA se constituera ». Une solution pourrait venir de l’IA elle-même : elle devrait être capable de répondre à nos questions pour nous aider à former un jugement éclairé. Mais un autre danger guette : “Information presented to humans will naturally inform (and bias) decision-making. Humans are quick to assume that algorithms or information provided by statistical analysis is somehow void of bias.”

EXTRAITS CHOISIS

« Il serait bien d’encadrer juridiquement cette notion vis-à-vis des entreprises et des organismes publics faisant des affaires au Québec. »

« C’est impossible. La seule chose à faire est d’établir ou de rétablir une confiance avec ceux qui construisent et sont propriétaires de ces technologies par un contrôle social et politique qui satisfasse le plus grand nombre des utilisateurs, en réduisant les abus et détournements que les créateurs et propriétaires de ces technologies pourraient être tentés de réaliser. »

« Probablement pas. Je crois qu’il est déjà impossible de donner un consentement éclairé pour des technologies informatiques qui ne se base même pas sur l’IA. Par

exemple, comment être sûr que nous ne sommes pas espionnés par les logiciels que nous achetons. »

“For decently complex systems, the user has to be fully made aware of how the data being generated can and might be used, along with theoretical guarantees or open code base proving their claims. But for very complex systems, here, even the creator wouldn’t know how the data might be used completely. But, even in the worst of the cases, rigorous proof of claims and possible benefits, analysis on a test group can help earn the trust of the user and allow the person to give consent.”

“As technology advances, the demands for our consent will increase exponentially. Under those conditions, the unaided human will not be able to give truly informed consent in many of the cases where it is demanded. The proof is that we already have become conditioned to signing off on agreements that we have not actually read or understood. The demands are only going to increase. The solution, if there is one, would involve “loyal” AI agents assisting us.”

5. FAUT-IL LIMITER L'AUTONOMIE DES SYSTÈMES INFORMATIQUES INTELLIGENTS? UN HUMAIN DEVRAIT-IL TOUJOURS AVOIR LA DÉCISION FINALE?

Beaucoup de réponses positives. L'être humain doit toujours être aux commandes, garder la main. L'IA est un outil, une aide à la décision. Les points de vue divergents sont toutefois intéressants : « L'IA est potentiellement plus précise, moins biaisée et bientôt plus créative que l'humain. Profitons-en ! » et dans le même ordre d'idée : « Les humains sont corruptibles. Une IA peut avoir un code moral plus strict que les humains. » Il pourrait aussi y avoir des contraintes liées au fait qu'on a besoin d'une prise de décision urgente.

Le contexte est évidemment important : faire des croissants ou lancer une attaque, ce n'est pas la même chose. L'humain devrait minimalement pouvoir prendre la décision d'arrêter un système autonome. Et cela ne semble pas négociable « dans le cas de décisions complexes incluant une dimension éthique engageant la responsabilité ».

EXTRAITS CHOISIS

« Il pourrait arriver un point dans le développement des systèmes où il sera possible de démontrer qu'un être humain n'a pas la capacité de prendre une meilleure décision que l'ordinateur. »

« La décision finale, non, car l'avantage de l'IA est de pouvoir prendre une décision instantanée en fonction d'une somme de paramètres qu'un humain ne pourrait jamais analyser aussi rapidement. Mais la responsabilité de la décision, elle, doit toujours être assumée par un humain. »

« Oui et oui, les systèmes informatiques sont des aides à la décision et doivent le rester. Pourquoi donner à un cyborg le pouvoir sur nous ? »

« Les décisions fondamentales doivent être humaines et reposer sur le plus large consensus possible. »

“You always want to have the option of an off switch. And we need to build systems in such a way that we can come to an understanding of how the machine is making the decision.”

“Obviously with the current state of the technology, you can’t let it have total control over too many things. That is unlikely to be true forever; eventually the AI is probably going to be smarter than the human... and possibly more benevolent than the human, which is where you should really be putting your energy. At some point the question may be whether the human should even get any input into certain decisions, especially into decisions that affected more than just that human.”

“If the human does NOT always make the final decision, then there needs to be a transparent interface so that users can correct the decision-making computer system when it makes mistakes (like google translate, you can provide a better translation).”

JUSTICE (ÉQUITÉ, SOLIDARITÉ, DIVERSITÉ)

1. COMMENT S'ASSURER QUE LES BÉNÉFICES DE L'IA SOIENT ACCESSIBLES À TOUS ?

Par un prix abordable (ou la gratuité), par l'*open source* et en exposant clairement quelles sont les décisions que l'IA prendra à notre place (transparence). Mais est-ce possible dans le système capitaliste que nous connaissons ? « Le secteur privé ne devrait pas pouvoir exploiter une rente à son seul profit et au détriment du reste de l'humanité. » On pourrait d'ailleurs taxer les compagnies qui s'enrichissent excessivement grâce à l'IA (cela nuirait-il à l'innovation ?).

L'éducation pourrait avoir son rôle à jouer pour lutter contre la fracture digitale. C'est le rôle des gouvernements (voire de l'ONU) que de répartir équitablement ces bénéfices et de s'assurer que les valeurs de l'IA soient alignées avec les valeurs humaines. Un *basic income*, un appel au réalisme politique tempère les attentes : « ne soyons pas utopistes, ce n'est pas l'IA qui crée les inégalités, c'est l'humain ». Un répondant note aussi que les technologies de l'information rendent possible une démocratie participative. Un autre évoque le *basic income*.

EXTRAITS CHOISIS

« Construire les IA dans l'intérêt commun plutôt que comme propriété privée. Réglementer pour forcer les formes avancées d'adopter une licence libre GNU par exemple et de privilégier le partage de l'information. »

« Il ne faut pas laisser l'IA entièrement à la merci de l'entreprise privée. »

« Les avancées possibles des AI, par exemple la découverte d'une nouvelle protéine, doivent être des biens collectifs. »

“General quality of life for everyone should be improved with AI. Legal system seems to be one that will be greatly affected and see a lot of change, for the better.”

« La fracture de l'accessibilité pourrait dépendre de la mainmise ou non de grands groupes privés sur les données générées par la population. »

« Il faut revoir en profondeur les lois internationales sur les brevets. Le développement des IA ne progressera véritablement que si l'information qui les sous-tend est du domaine public. L'appropriation de cette technologie par des groupes d'intérêts spécialisés (corporations, armée, gouvernements) ne doit pas être rendue possible, sans quoi elle sera inévitablement détournée pour servir ces intérêts plutôt que les citoyens. »

« Faire de l'équité un chantier central. Inclure les chercheurs et les groupes communautaires qui exercent la collaboration dans le design de solutions équitables. Voir les travaux de l'Unité soutien (SRAP) et du chantier Mobilisation et participation citoyennes d'Alliance santé Québec. »

“Give free Wi-Fi to the poor for starters.”

“This is a very complex question. One could argue that everyone already benefits from AI through “free” products like Facebook and Google Maps. What is missing is an understanding of the market value of someone’s data relative to the machine’s ability to build a more powerful model. Governments at all levels need to be using AI with the data they currently manage as another part of their policy-making tool set.”

2. FAUT-IL LUTTER CONTRE LA CONCENTRATION DU POUVOIR ET DE LA RICHESSE AU SEIN D'UN PETIT NOMBRE D'ENTREPRISES EN IA ?

Les réponses sont nettement positives. En favorisant l'*open source* et les licences libres GNU. Car c'est l'État plutôt que le secteur privé (les GAFA) qui a la confiance des citoyens. Les inquiétudes sont réelles : « La démocratie survivrait-elle avec une IA prédominante dans de mauvaises mains ? » Comment faire, toutefois ? On n'est même pas arrivé à faire que le logiciel libre supplante le logiciel propriétaire. Nationaliser pour rester « maître chez soi » ? Quoi qu'il en soit, l'IA devrait être vue comme un bien commun qui ne sert pas une minorité. Un répondant évoque la nécessité d'un organisme antitrust pour briser certains monopoles.

Cependant certains valorisent un modèle plus concurrentiel : « Si certaines entreprises parviennent à se créer une niche qui leur rapporte pouvoir et richesses, grand bien leur fasse. Mais la connaissance doit être du domaine public afin de favoriser la concurrence. » Un répondant propose de faire des données personnelles la propriété des individus qui pourraient se prévaloir également d'une IA d'assistance personnelle loyale envers eux.

« Évidemment, il faut lutter contre la concentration du pouvoir, point. »

“Yes, it seems there will be a lot of power available to those who control AI systems. New legislation/law will be required to monitor this, along with taxes on automation, etc.”

« Il faudrait surtout que les programmes de base soient universels et bâtis pour le bien commun. Sinon ce ne seront que des robots au profit de ceux qui dirigent déjà malicieusement le monde dans leur propre intérêt, alors soit ça ne changera rien, soit ça empirera les inégalités, la violence, les conflits, etc. »

“The hands of a small number of AI companies or the hands of a small number human entities (i.e. the 1%) should not have more power and wealth than the 99% of human beings on earth. Powerful entities should adopt socially responsible behaviours at all time, especially when in presence of the public. (...) The democratization of AI should definitely empower the 99% of human beings.”

3. QUELLES SONT LES DISCRIMINATIONS QUE L'IA POURRAIT CRÉER OU EXACERBER ?

Toutes les formes de discrimination « classiques » semblent pouvoir être exacerbées par l'IA, en particulier les discriminations « sociales, raciales, économiques », mais aussi « linguistiques et culturelles ». Entre les gens, mais aussi entre les groupes ou entre les États. Un scénario dystopique se profile : celui où une nouvelle classe d'ultra-riches (le 1% ?) utilise l'IA pour perpétuer les inégalités socio-économiques. On mentionne aussi que l'accès à la technologie peut être exclusif et excluant, en particulier pour les personnes plus âgées.

Un répondant précise le type de mécanisme que pourrait encourager l'IA : « L'IA peut être le bouc émissaire parfait sous forme d'une BLACK BOX : Pourquoi je n'ai pas reçu une marge de crédit M. le directeur de la banque ? Ah, c'est le système qui nous a donné le résultat, je suis désolé. » Il semble aussi clair pour les répondants que ce sont les humains en tant qu'individus ou en tant que groupes (ex. racisme systémique) qui sont et seront responsables de ces discriminations — pas l'IA.

EXTRAITS CHOISIS

« [Il faut se méfier de l'] apparition d'une "caste" des experts en IA, connus ou occultes, détenant le savoir, donc le pouvoir. [Il faut aussi se méfier des] discriminations selon l'état de santé (flirt avec l'eugénisme), discriminations raciales, sexuelles, envers les gens âgés, envers les femmes, etc. [Attention enfin aux] discriminations économiques, accentuant la pauvreté du plus grand nombre et le pouvoir des riches sur les décideurs. »

« Il y en a trop... Là est justement le problème. Nous avons de la difficulté à établir ce qu'est

une discrimination ou si nous en faisons déjà. Comment l'IA pourrait le déterminer à notre place sans justement utiliser les mêmes discriminations que nous lui fournissons en lui donnant des datas. »

« Les réseaux sociaux sont déjà source de stéréotypes, propos racistes, sexistes, stigmatisants. On peut penser à filtrer cela, ce que déplace le problème. Ces filtres aussi pourraient comporter des discriminations indues. »

« Les algorithmes devraient être développés par des équipes multidisciplinaires et multiculturelles afin de ne pas perpétuer de préjugés de genre, d'économie, d'ethnie, etc. »

« Si l'IA participe au bien-être et à l'autonomie, les personnes qui en ont besoin, mais n'y ont pas accès, seront d'autant moins bien loties. »

« Si l'IA est déployée par des groupes d'intérêts spécialisés (armées, gouvernements, corporations), elle ne servira que leurs intérêts du moment au détriment de la population. »

"See weapons of math destruction. AI models with labelled training data that is discriminatory will simply perpetuate and reinforce these discriminations."

*"It's going to be hard to deal with that, because in order to admit that the AI is going to find a regularity, you have to admit that the regularity exists. You have to be willing to say, "Yes, XXX people *are* more likely to default on loans, but we want to ignore that anyway". After that, it's a relatively simple technical problem to make the AI implement your wishes. Short-term AI, anyhow."*

4. LE DÉVELOPPEMENT DE L'IA DEVRAIT-IL ÊTRE NEUTRE OU CHERCHER À RÉDUIRE LES INÉGALITÉS ÉCONOMIQUES ET SOCIALES ?

La plupart des répondants sont favorables à une IA qui contribuerait activement à réduire les inégalités économiques et sociales. Plusieurs y voient même une priorité. Un répondant optimiste pense que les réductions des inégalités seront un effet mécanique du développement de l'IA. Un autre voudrait qu'elle promeuve surtout l'égalité des chances. Toutefois, quelques sceptiques préféreraient qu'elle reste neutre : « Qui va indiquer quelles inégalités réduire ? » Et les plus pessimistes soutiennent qu'il y aura toujours des inégalités... ce qui n'empêche pas d'essayer de les réduire. Enfin, un répondant suggère que l'IA devrait rester neutre quant aux inégalités économiques et sociales lorsqu'il s'agit d'usages commerciaux, mais que les usages non commerciaux devraient viser davantage d'égalité.

EXTRAITS CHOISIS

« Oui l'intention devrait toujours être celle-ci, mais aussi la réduction des impacts environnementaux. »

« L'IA ne peut être neutre, donc autant assumer une direction meilleure pour tous. »

« Pourquoi développons-nous l'IA ? La réduction des inégalités ne me semble pas être la raison première ; cela ne veut pas dire cependant que le développement de l'IA devrait être neutre : les inégalités économiques et sociales pourraient faire office de "site contraint", pour que le développement ne se fasse pas au détriment de valeurs importantes. »

"AI models should be applied within a policy framework. No information system is neutral and any architect or policymakers must embrace the ethical challenges and opportunities when applying AI. In this context, reducing existing inequalities is a moral imperative. Machine learning models need to be conceived inside of a larger pipeline that can mitigate regressions and provides recourse for error."

"But we should make sure that by doing so we are not actively causing friction between different groups or trying to homogenize them. The effect, in that manner, should be neutral."

"It should be neutral in commercial settings, otherwise the technology might never be adopted at all—leading to no benefit to the society. But it should also reduce socioeconomic inequalities in a non-commercial setting by giving

everyone access to the same tools and opportunities."

5. QUELS TYPES DE DÉCISIONS DE JUSTICE POURRAIT-ON DÉLÉGUER À UNE IA ?

Il ressort des contributions qu'aucune décision importante ne devrait être déléguée à une IA. L'IA ne doit être qu'un outil d'aide à la décision. Elle pourrait ainsi « accélérer le traitement des dossiers », voire « prendre des décisions faciles après une analyse des preuves », comme des décisions liées au paiement des contraventions.

L'IA pourrait être bénéfique dans d'autres aspects de la justice : « Détecter un mensonge ou un faux souvenir. Détecter les risques de récidives. » Si une intelligence artificielle générale était développée, alors on pourrait envisager que des IA remplacent des juges ; mais cette option est très loin de faire consensus, même s'il est avéré que les juges humains sont souvent biaisés dans leurs jugements et soumis à diverses pressions. Peut-être faudra-t-il un jour repenser l'institution judiciaire de fond en comble pour rendre possibles des « jugements artificiels ». Quoi qu'il en soit, la réduction des coûts et la démocratisation de la justice seraient une bonne nouvelle et l'IA pourrait certainement y contribuer, par exemple en facilitant l'accès à la jurisprudence.

EXTRAITS CHOISIS

« L'IA pourrait remplacer la personne qui prend des notes. »

« L'IA serait plus juste parce que non soumise aux émotions ou à la pression des médias et autres groupes d'opinion et de pression. La seule chose qui serait éventuellement à revoir serait le Code pénal, compte tenu des différences observées entre un

jugement humain et artificiel. »

“I don’t believe any final decisions should be made by the AI. Seems the legal aid/technician and data processing could be best managed by AI.”

« Les décisions impliquant le recours au jugement pratique complexe (jurisprudence) devraient être réservées aux humains. La justice est aussi un processus social. Ne l’oublions pas. »

« L’IA pourrait servir de recherchiste à la population (ainsi qu’aux juristes), en ayant accès à toute la jurisprudence. Ce travail démocratiserait l’accès à la justice puisque l’essentiel des coûts assumés par les citoyens est relié au temps passé par les juristes à faire ces recherches. »

“Current and near-future AI aren’t going to be able to comprehend the law or apply it other than in cases so mechanical that you don’t really need “AI” at all. I suspect that any real legal decisions will take a truly general intelligence.”

“AI predictive technology can be used to help judges make better decisions. The idea is not to replace judges.”

VIE PRIVÉE (INTIMITÉ)

1. COMMENT L’IA PEUT-ELLE GARANTIR LE RESPECT DE LA VIE PRIVÉE ?

Plusieurs répondants s’interrogent sur la pertinence de cette question : Comment l’IA peut-elle garantir ce respect ? L’impression est plutôt qu’elle la viole, à répétition et sans le consentement des utilisateurs. Il semble même y avoir une contradiction : l’IA a besoin de nos données pour se développer.

Mais il existe peut-être des options : « crypter tout », ne pas être invasif dans la demande de données personnelles. Quelqu’un remarque : « Le respect de la vie privée est garanti si la personne n’est pas exposée à une IA par défaut. » Il en va aussi de la responsabilité des utilisateurs : « C’est à chacun d’entre nous de contrôler son exposition : allez faire vos courses dans des boutiques indépendantes et payez en liquide, plutôt que d’acheter sur internet. »

On retrouve aussi une méfiance vis-à-vis du secteur privé : « Rien n’est garanti si c’est géré uniquement par le privé ». C’est pourquoi on appelle l’État et le législateur à la rescousse : il faudrait que les lois québécoises relatives à la vie privée soient respectées et améliorées. « C’est un gros défi », car ne serait-il pas déjà trop tard ? Nos données Facebook, par exemple, ont peut-être été siphonnées depuis longtemps par Cambridge Analytica ou une compagnie équivalente. Et c’est sans parler des « hackers ».

EXTRAITS CHOISIS

« Je crois que l’économie de l’information, basée sur la traçabilité, peut occasionner plus de partage d’information, mais en même temps plus de transparence dans leur usage, et donc avoir ses informations partagées n’aura pas de conséquences aussi lourdes si ceux qui la visionnent sont tracés aussi. »

“Let’s face facts, there are, realistically speaking, no truly reliable guarantees that AI can respect people’s privacy. Health records & private accounts are hacked all the time despite the best security upgrades that technology has to offer. Google reads our private e-mails, doesn’t it?”

“Differential privacy—the idea that you can give away information about yourself without ever having it trace back to you as the source. But, if such a practice is possible and can be made prevalent then I believe that informed consent is possible. The user has to be fully made aware of how the data being generated can and might be used, along with theoretical guarantees or open code base proving their claims.”

“Make people’s private data truly their private property.”

2. NOS DONNÉES PERSONNELLES NOUS APPARTIENNENT-ELLES ET DEVRAIT-ON AVOIR LE DROIT DE LES EFFACER ?

Les réponses sont massivement positives pour les deux sous-questions. Quelqu’un précise « et cela devrait être très facile comme procédure, pour que tous puissent le faire. » Un répondant s’oppose à l’idée que nos données nous appartiennent, mais cela n’empêche pas que nous devrions avoir « un droit de regard sur leur utilisation. » Si la majorité des répondants admet implicitement que ce sont les individus qui devraient posséder leurs données, certains l’envisagent plutôt comme un bien collectif.

L’effacement des données ne devrait toutefois pas entraver la justice (ou les services de santé) qui pourrait avoir besoin d’avoir accès à des données anciennes. Cet effacement ne devrait pas non plus causer de tort à autrui.

EXTRAITS CHOISIS

« Oui, chaque citoyen devrait être propriétaire de ses données privées au même titre que les artistes de leurs productions culturelles. »

« Non, mais les données devraient être considérées comme un bien national, comme les bibliothèques ou les réserves naturelles. »

« Absolument et de façon non équivoque. Seules les données essentielles pour le bon fonctionnement du gouvernement devraient être conservées : démographie, revenu, santé, judiciaire. Toutes les autres devraient pouvoir être contrôlées par l’utilisateur. »

“As long as companies own and licence IP, individuals should have a right to all data they create.”

“Generally yes. But I have a very broad definition of what should be considered personal data (and should be private property). Within this larger view even our criminal records would be personal data that we own (though not without controls). It would be a category of personal data that we should not be able to delete—at least not whenever we choose.”

3. DEVRAIT-ON SAVOIR À QUI NOS DONNÉES PERSONNELLES SONT TRANSMISES ET, PLUS GÉNÉRALEMENT, QUI LES UTILISE ?

Oui, répondent les gens à l'unanimité ! Quelqu'un précise : « Tout comme nous devons savoir qui entre dans notre maison, nous devons savoir qui accède à nos données personnelles. » Un autre : « Oui, [et on devrait savoir] à qui, comment et dans quel but ». Un répondant remarque qu'on risque de se lasser de savoir qui utilise nos données et qu'on pourrait assez vite s'en désintéresser. Mais cela n'empêche évidemment pas qu'on ait le droit de le savoir.

EXTRAITS CHOISIS

« Oui, je pense que je devrais même avoir un portail où je contrôle 100 % de la donnée que je donne. »

« Nos données ne devraient jamais être transmises sans qu'au préalable une demande claire et concise ait été faite en ce sens. Pas de contrat de 20 pages en petites lettres où l'on doit deviner qu'il y a là une permission donnée *ad vitam aeternam*. Si l'on s'abonne à un service, l'information ne devrait jamais pouvoir être utilisée autrement que pour le service demandé. »

"Absolutely and they should be required to ask permission to do so on a regular basis. Permission is not granted in perpetuity."

"Absolutely. Personal data should be private property. We should defend it and allow the owner to control who can access it and to what extent they can access it. The current default—wherein we cede our data

to others—is bad for citizens and bad for democracy. There is another option."

4. EST-IL CONTRAIRE AUX RÈGLES D'ÉTHIQUE OU D'ÉTIQUETTE QU'UNE IA RÉPONDE À VOTRE PLACE À VOS COURRIELS ?

Cette question soulève des intuitions contradictoires. Plusieurs remarquent que ce type de service existe déjà ou que certaines personnes ont des assistants humains qui répondent à leur place à leurs courriels. Une option serait que l'IA prépare la réponse, mais que celle-ci soit validée par l'humain (qui aurait ainsi le « dernier mot »). Un répondant précise que « l'important selon moi est que la personne qui utilise ce service ait la confiance et la compréhension nécessaires du service. » Une autre demande à ce que le procédé soit transparent, c'est-à-dire que l'interlocuteur sache que la réponse à son courriel provient d'une IA. Il n'y a peut-être pas de réponse générique à cette question : ça dépend des types de questions (« Es-tu disponible pour ce RV ? » vs « Penses-tu qu'on devrait embaucher telle personne ? »).

EXTRAITS CHOISIS

« Non du moment qu'il est stipulé de façon lisible que la réponse a été produite par une IA en lieu et place de l'utilisateur concerné. Si ce dernier veut utiliser une IA pour répondre à sa place, il en est de sa responsabilité... du moment que le fournisseur de service internet offre la possibilité d'activer ou de désactiver cette fonction. Il est bien entendu qu'il ne s'agit pas d'imposer un tel service. »

« Utile pour ceux qui ont à gérer un grand volume de messages similaires et peu complexes. »

« Ça dépend, si vous répondez toujours la même chose pour la même question, ça ne ferait pas de différence pour vous. »

« S'il est question de service à la clientèle, de répondre à un besoin humain à satisfaire qui engage la responsabilité d'autrui, je m'attends à ce que ce soit un humain qui réponde. »

"Yes. Human intent is a critical component to our society's framework. We can delegate to AI, but human dignity demands that you should know if you are interacting with a machine."

"Similarly, if an organization has a bot deal with people, it should always identify itself as a bot. People should always know if they are dealing with a human or a machine. And the organization that has bots dealing with people should always be held responsible for any actions the bot takes on the organization's behalf."

5. QU'EST-CE QU'UNE IA POURRAIT FAIRE EN VOTRE NOM ?

Une question ouverte qui suscite des réponses très diverses, allant de « rien » à « tout » (pour autant que l'on y a consenti). Entre les deux : programmer un rendez-vous, gérer mes finances, mon agenda, faire mes impôts et autres tâches administratives, voter (!). Mais je devrais toujours être tenu responsable des conséquences de ce que l'IA fait en mon nom. (Plusieurs répondants confondent cette question avec « Qu'est-ce que l'IA pourrait faire pour vous ? », par exemple : passer l'aspirateur).

EXTRAITS CHOISIS

« Tout ce que j'aurais préalablement approuvé. »

« Toutes tâches qui n'engagent pas une décision pour l'avenir. »

« Rien de sérieux pouvant avoir des implications légales ou émotionnelles. »

"Book appointments respond with numerical data that is already in the public domain, check on the well-being of family pets."

*"That depends on the AI. I wouldn't trust any *present* AI to do anything that I couldn't countermand or that people would interpret as a direct application of my personal judgment."*

"My recommendation is to adopt a paradigm in which each citizen owns private, 'loyal' AI tools (agent) that can help protect, manage, analyze and use a citizen's private data (stored in a protected online profile) to help that citizen at their

behest and only their behest. (...) Some people might say they can do simple repetitive tasks, perhaps review email. Others might allow their AI agent to browse the web to plan online shopping. Others might let the agent actually make purchases autonomously. Others might allow the AI agent to perform investment transactions for them. In an advanced future, some prefer to trust their AI to participate in a family vote about 'pulling the plug', given on its intimate access to its owner's private data, which could analyzing a variety information types taken from a personal profile, allowing it to use predictive analysis to help decide what the citizen might want if they were able to speak."

CONNAISSANCE (PUBLICITÉ, PRUDENCE)

1. LE DÉVELOPPEMENT DE L'IA FAIT-IL COURIR UN RISQUE À LA PENSÉE CRITIQUE ?

Les réponses sont contrastées, mais penchent plutôt pour le non. Du côté du « risque », on craint plusieurs choses : une perte de curiosité, la publicité, une normalisation de la pensée et la mise à l'écart des points de vue marginaux. Il se pourrait aussi que l'IA parle au nom des humains et qu'elle paraisse trop fiable : « La machine ne peut se tromper ; tout est dit ; il n'y a plus rien à ajouter. »

Du côté des avantages, plusieurs notent que le temps gagné par l'automatisation pourrait être investi dans la pensée critique et le fait que l'IA

et les technologies de l'information rendent plus accessible l'information, voire qu'on pourrait programmer l'IA pour avoir une pensée critique — on perçoit aussi cette idée que l'IA pourrait être plus neutre que des humains. Enfin, on peut voir l'émergence de l'IA comme une belle occasion — ou une nécessité — pour les humains d'exercer leur pensée critique.

EXTRAITS CHOISIS

« Oui, mais pas si elle s'applique à faciliter la vie des gens en leur laissant plus de temps pour s'instruire et donc développer leur pensée critique. »

« Non, au contraire. La somme des connaissances humaines croît de façon exponentielle, au point qu'il devient impossible de connaître tous les tenants et aboutissants d'un problème. L'IA, avec ses capacités de synthèse, permet aux humains de filtrer l'information redondante et de se concentrer sur l'essentiel. »

"I believe it certainly could compromise humans quest for knowledge & need to problem solve & therefore seriously impair our critical thinking & problem solving capacities & increase depression in people who may in future, have no motivation to use their god-given gifts & intelligence because they have been replaced by AI."

"It would definitely be more of a crutch than a tool if we become overly reliant on it. Instead the

development and the products that are created using AI tech should be such that it aids critical thinking, aids skill development and indirectly making life easier.”

2. COMMENT MINIMISER LA CIRCULATION DE FAUSSES NOUVELLES OU D'INFORMATIONS MENSONGÈRES ?

Une question ouverte qui génère des pistes de solutions très diverses : soutenir financièrement les médias (locaux, traditionnels) qui vérifient l'information, investir dans le journalisme de qualité (qui multiplie les sources d'infos), éduquer les gens, utiliser une IA pour vérifier une info, punir ceux qui mettent de fausses nouvelles en circulation, les effacer, imposer des règles aux plateformes (type Facebook) qui font circuler ces fausses nouvelles. Notre accoutumance collective aux nouvelles « gratuites » (en un certain sens seulement) est aussi pointée du doigt.

Faut-il censurer les fausses nouvelles ? Un répondant prend position : « Il faut plutôt diffuser au maximum des articles de vérification des nouvelles, car la censure est contre-productive (elle peut par exemple alimenter des théories du complot) ». Un point de vue pessimiste : “It may become impossible as AI advances so too will its ability to mimic voices and fabricate images and video.”

EXTRAITS CHOISIS

« Redéfinir le métier de journaliste. Développer un système d'accréditation des sources d'information. Reconnaître des experts en communication dans les différents secteurs de l'activité humaine. »

« Il y aura toujours de fausses nouvelles, il faut développer l'esprit

critique et éduquer les jeunes en ce sens. »

« Il ne faut pas que la censure provienne directement de l'IA, par contre l'IA peut devenir un outil qui permet de prédire la probabilité pour qu'une nouvelle soit fausse. »

« Éduquer les gens à la pensée critique, à la recherche d'information crédible et à l'ouverture de leur conscience. »

3. LES RÉSULTATS DES RECHERCHES (POSITIFS OU NÉGATIFS) EN IA DOIVENT-ILS ÊTRE DISPONIBLES ET ACCESSIBLES ?

La réponse est sans ambiguïté positive. Et ce devrait être le cas pour tous les résultats de recherche dans tous les domaines, soutiennent plusieurs répondants. Ces résultats, précisent d'autres répondants, devraient être *open source* (notons qu'ils le sont déjà dans une très large mesure).

EXTRAITS CHOISIS

« Tout à fait. Et le plus possible, vulgariser ces résultats pour les rendre accessibles à tous. Pas de résultats opaques, avec des termes incompréhensibles... »

“This question has more to do with research than AI. Publicly funded research, with few exceptions, should be made available as a Social Good.”

“Yes. I know people who think really powerful results should be kept from the “bad guys”. That is a total pipe

dream. All you'll do by trying is to disadvantage the "good guys". Your best bet is to be open."

"YES!! Especially negative results. They would provide as much information, if not more about a particular problem."

4. EST-IL ACCEPTABLE DE NE PAS ÊTRE INFORMÉ QUE DES CONSEILS MÉDICAUX OU LÉGAUX SONT DONNÉS PAR UN « CHATBOT » ?

Pour les répondants au questionnaire, c'est largement le non qui l'emporte. Deux préoccupations semblent guider ces réponses ; le souci de transparence et celui de prudence : « Les conseils donnés par le chatbot peuvent être pris en considération de façon différente si la personne sait si elle parle avec un humain ou un chatbot. Un chatbot ne peut connaître toutes les variables d'une situation. » Plusieurs remarquent qu'il est d'ailleurs facile d'informer une personne qu'elle communique avec un chatbot.

EXTRAITS CHOISIS

« Éventuellement oui. Aucun passager d'avion ne demande au maître de cabine si c'est le pilote ou l'autopilote qui contrôle l'avion. »

"The source of such advice being often critical to a person's well-being, one should be aware of the source of this information."

"No, every information should be presented along with the source exactly as it is, along with the analysis of how accurate or biased the information/advice might be. It may happen that the person may

rely on that information even after realizing that it is from a chatbot, as it would get good results. And that is the kind of relationship we'd like to foster."

5. EN QUEL SENS LES ALGORITHMES DEVRAIENT-ILS ÊTRE TRANSPARENTS QUANT À LEUR PROCESSUS DE DÉCISION ?

Cette question laisse beaucoup de répondants dubitatifs. La réponse qui revient le plus est « le plus possible » en ayant conscience des difficultés techniques en jeu ici (c'est-à-dire du problème de la « black box »). Si certains pensent que l'IA ne devrait tout simplement pas prendre de décision, les autres semblent d'accord pour qu'une IA prenne une décision, à condition d'avoir accès à une « justification déchiffrable par un humain ». Il se pourrait aussi que dans certains contextes, la transparence ne soit pas désirable. Plusieurs remarquent que la transparence sera importante pour susciter la confiance envers l'IA. Un répondant suggère de donner le degré de fiabilité d'une décision prise par une IA.

On note aussi que la transparence implique de savoir à partir de quelles données (ou type de donnée) une IA prend une décision et les valeurs (ou intérêts) qui guident sa décision.

Un participant suggère au contraire qu'on ne devrait pas être plus exigeant envers une IA qu'envers un humain.

EXTRAITS CHOISIS

« Une description du processus de décision des algorithmes devrait venir avec l'achat d'un produit IA, comme les modes d'emploi ou les garanties du fabricant qui accompagnent un produit lors de son achat. »

« Si les créateurs d'une IA ne sont pas capables de définir précisément la portée et les limites de capacité de décision d'une IA qu'ils proposent, celle-ci ne devrait pas pouvoir être commercialisée. »

« Les échelles de valeurs utilisées pour leur prise de décision. Voir les valeurs relatives des différents éléments décisionnels. Par exemple : Chat vs chien, collectif vs individu, etc. »

« Totalelement transparent. Comment faire confiance si on ne sait pas sur quels principes ils se basent pour faire leur analyse ? Tout comme il est toujours pertinent de comprendre la méthodologie de recherche utilisée par des chercheurs. »

"You should be able to ask an AI why it made a choice then if you find its reasons lacking you should be able to make it change its behaviour."

"We may be able to infer decision-making processes but we should not assume that there is any internal motive or intent in an algorithm."

DÉMOCRATIE (PUBLICITÉ, DIVERSITÉ)

1. Faut-il contrôler institutionnellement la recherche et les applications de l'IA ?

La réponse est globalement positive, en particulier pour les applications de l'IA (la liberté de la recherche scientifique est une valeur importante). On suggère un « bureau de l'Ombudsman de l'IA » et des comités d'éthiques de l'IA ou encore une sorte de serment d'Hippocrate. On note aussi que « le sujet est trop éminemment politique et social pour être laissé aux mains du privé ». Ce contrôle, toutefois, ne devrait pas nuire à l'innovation (pour autant que celle-ci soit compatible avec le bien commun et les droits humains). Une difficulté inhérente à ce contrôle institutionnel relève de la politique internationale : comment des pays aux intérêts divergents pourront-ils se mettre d'accord sur des institutions communes ?

EXTRAITS CHOISIS

« Oui, à condition que nous développons une démocratie participative et que les gouvernements soient d'abord au service de la majorité, pas à celui du capital. »

« Non, mais poser des bornes est essentiel. »

"Yes but good luck getting China or Russia to follow along."

"Controlling AI research is simply not possible. The research itself should continue, but a broader communication framework explaining what AI can and cannot do is critical. Sensitizing researchers to the ethical ramifications of their work is also important (e.g. the Hippocratic oath)."

2. DANS QUELS DOMAINES EST-CE LE PLUS PERTINENT ?

Question ouverte. Beaucoup répondent « dans tous les domaines ». La santé arrive largement en tête des domaines les plus souvent cités. On trouve ensuite (dans l'ordre) : l'armement, la justice, l'environnement, l'alimentation, la surveillance, la vie privée, la finance, la sécurité, l'éducation, le gouvernement. Sont aussi mentionnés : l'économie, l'industrie, l'épigénétique, le journalisme, le transport, les services municipaux, la recherche sur une super-IA (AGI), les voitures autonomes et les publicités ciblées.

EXTRAITS CHOISIS

« Dans tous les domaines liés à la vie (biologie) et à la vie en société. »

3. QUI DEVRAIT DÉCIDER — ET SELON QUELLES MODALITÉS — DES NORMES ET VALEURS MORALES DÉTERMINANT CE CONTRÔLE ?

Les répondants, souvent très indécis sur cette question, hésitent entre diverses options : le parlement, des consultations publiques, l'ensemble de la population (référendum, tirage au sort), un comité multidisciplinaire (experts, élus, citoyens), la commission d'éthique en sciences et technologie, un « comité de sages », une institution internationale (type ONU). L'idée que cette instance de décision devrait être indépendante (du pouvoir politique et économique) ressort à plusieurs reprises. On trouve aussi le souci que cette instance soit représentative de la diversité des citoyens.

EXTRAITS CHOISIS

« Je ne sais... Un comité conjoint, multidisciplinaire, universitaire, populaire et impartial. »

« Nous tous, en développant des modes d'information, de

consultations et de prises de décisions qui impliquent le maximum de gens de toutes provenances. Pas la "démocratie" actuelle. »

« Bien des comités. Ceux-ci pourraient établir des règles, valeurs, etc., en lien avec chaque institution où il y aurait l'un de ces comités. Ainsi, ceux-ci pourraient établir une sorte de « charte », que devrait suivre l'institution et faire des recommandations... Qui bien sûr ne devraient pas être tablettées !

« Au Québec, la Commission d'éthique en science et technologie a déjà produit un document sur les villes intelligentes en indiquant les enjeux à considérer. D'autres projets IA pourraient être analysés par cette instance ou par d'autres instances gouvernementales spécialisées dans le domaine visé. Un ombudsman pourrait être nommé pour certifier des projets IA et recevoir des signalements liés au non-respect des principes de la Déclaration de Montréal en IA. »

« Comme l'IA touchera tous les domaines (loi, santé, science, société, arts), il est vital que des spécialistes de chacun de ces domaines soient représentés au sein de l'organisme. Le gouvernement doit financer

adéquatement l'organisme, mais ne peut intervenir dans son fonctionnement. De plus, un gouvernement ne devrait pas avoir le pouvoir d'abolir l'organisme ou d'entraver son travail. »

"Canadian's from all groups, backgrounds & beliefs."

"This should function like an IRB as in the drug development and testing industry."

4. QUI DEVRAIT CHOISIR LES « RÉGLAGES MORAUX » DES VOITURES AUTONOMES ?

Cette question suscite des réponses très variées : le parlement, une agence gouvernementale, l'État, les pouvoirs provinciaux, l'État en collaboration avec l'industrie, un comité d'expert en éthique, la SAAQ, le manufacturier qui porte la responsabilité de l'auto, une autorité de certification des logiciels, un comité d'usagers, les juges de la Cour suprême, un organisme international type ONU. L'utilisateur pourrait aussi avoir le choix de certaines options. En passant, on constate aussi chez plusieurs répondants une défiance envers les voitures autonomes (« elles devraient être interdites »).

EXTRAITS CHOISIS

« Encore là, ça pourrait être des comités de citoyens. Il faudrait un représentant des piétons, un autre des personnes âgées, un autre des moins de 16 ans, un autre pour les vélos, etc. Chacun pourrait s'exprimer sur les réglages moraux à donner aux voitures autonomes. »

« Sûrement pas les compagnies qui les construisent ! »

« Un commissaire en éthique et le Bureau du Coroner au Québec. »

"It should be a multilateral decision (after thorough public discussion)."

"Judges/supreme court, whoever decides and uphold the existing ethical guidelines should have a major role to play in the decision. But along with them, community participation, transport businesses and authorities, AI researchers and developers."

5. FAUDRAIT-IL DÉVELOPPER UN OU DES « LABELS ÉTHIQUES » POUR LES IA, LES SITES WEB OU LES ENTREPRISES QUI RESPECTENT CERTAINS STANDARDS ?

Une forte majorité répond par l'affirmative, ce serait une « bonne idée », « un bon début ». Cela pourrait ressembler à une norme ISO. Un répondant se demande toutefois pourquoi toutes les entreprises et les sites web ne devraient pas respecter ces standards. Un autre précise : « oui au cas par cas avec une charte standard ». Cela soulève aussi un certain scepticisme : Ces certifications seront-elles respectées ? Ne risquent-elles pas d'être corrompues ?

EXTRAITS CHOISIS

« Des certifications sujettes éventuellement à révision de façon vigilante pour s'adapter à telle ou telle situation. »

"Communities are different, people are different. (...) We should make

sure that by doing so we are not actively causing friction between different groups or trying to homogenize them. The effect should be neutral.”

“Definitely, at least three majors should be developed: corporate, government and individual ethical labels.”

RESPONSABILITÉ (PRUDENCE)

1. QUI SONT LES ACTEURS RESPONSABLES DES CONSÉQUENCES DU DÉVELOPPEMENT DE L'IA ?

Les répondants identifient plusieurs acteurs responsables : les universités, les chercheurs, les entreprises, les éthiciens, les politiciens, ceux qui commercialisent les applications, le gouvernement, les décideurs économiques, ceux qui en tirent un bénéfice financier, les représentants élus, la société, les utilisateurs, chacun de nous. Mais c'est peut-être « les développeurs/créateurs, les entreprises et le gouvernement » qui reviennent le plus souvent. Certains font le parallèle avec les animaux domestiques ou les enfants : ce sont leurs propriétaires/tuteurs/parents qui sont responsables. Dans le cas des IA, ce pourraient être les propriétaires. On évoque aussi ceux qui testent les IA, qui autorisent leur déploiement.

EXTRAITS CHOISIS

« Les gens qui les fabriquent, les gens qui les distribuent, et si on peut les pincer, les gens qui les utilisent de façon malveillante pour nuire, blesser, tuer ou dominer les autres (incluant les animaux) ou dégrader l'environnement. »

« Tous les membres de la chaîne

d'approvisionnement : du chercheur gradué, à la firme multinationale, en passant par les organisations de réglementation nationales, régionales et locales. »

« Les compagnies qui offriront les services devront être imputables et responsables, mais surtout les dirigeants de ces compagnies.»

“Whoever provides the results/ predictions of the AI decision-making. For example, Google is responsible for Google Translate.”

“Researchers developing models are partially responsible, however the application of AI ultimately rests with the owner/operators.”

2. COMMENT DÉFINIR UN DÉVELOPPEMENT PROGRESSISTE OU CONSERVATEUR DE L'IA ?

Une question qui reste souvent en suspens. Le développement progressiste rime avec collectif, transparence, moins d'écart de richesses. Le développement conservateur est assimilé à une certaine prudence : il ne faut pas se précipiter, il faut y aller graduellement. Quelqu'un remarque qu'il semble plus facile d'adapter la législation à l'IA que l'IA à la législation parce que le progrès va vite et qu'il semble difficile à freiner. Un autre que le développement progressiste devrait « favoriser les recherches alternatives ». Et un sentiment partagé par plusieurs : « On y va, on y va, peut-on faire autrement ? »

EXTRAITS CHOISIS

« En faisant des forums sur le sujet ! :-) Plus on en discutera,

de façon inclusive, plus il y aura un développement progressif de l'IA, dans le bon sens. Et aussi par l'éducation. Plus notre société sera éduquée, plus elle sera informée et fera des choix éclairés. »

« Pour le bien commun vs pour le bien privé. »

"It is progressive when it is maximizing freedom and agency. It is conservative when it is carefully monitored and cultivated as to insure safety."

"Conservative development: Checking, testing at each and every step. First in isolation, then within an isolated test group, and gradually deploy the AI."

3. COMMENT RÉAGIR DEVANT LES CONSÉQUENCES PRÉVISIBLES SUR LE MARCHÉ DU TRAVAIL ?

Plusieurs idées reviennent : un filet social solide ou revenu universel (*basic income*), une réforme de la fiscalité avec une taxe sur les robots ou une meilleure répartition des richesses. C'est cependant le recours à l'éducation et de la formation qui est le plus souvent préconisé : les gens devront s'adapter, ce qui demandera plus de formation continue. La transition devra certainement se faire graduellement et être transparente : les gens devront être tenus informés. Mais tous ne sont pas inquiets : « Le marché du travail a toujours été en évolution et continuera ainsi ». Par ailleurs, plusieurs semblent espérer une libération du travail.

EXTRAITS CHOISIS

« Offrir un salaire garanti en échange de participation à l'information des communs [numériques]. »

« Le travail n'est pas l'horizon de l'humanité ni son but. Le temps libre dégagé et les gains de productivité générés devraient être mis en commun pour permettre à tous de moins travailler sans perdre en niveau de vie. »

« En redirigeant les personnes vers d'autres types d'emploi qui participeront à plus de cohésion sociale. »

« Besoin de ralentir la cadence ; il faudrait se donner des priorités qui visent d'abord à développer ce qui va être au service de l'humain avant ce qui va remplacer l'humain. »

"AI tax, job displacement compensation, basic living wage, and research/development of new jobs."

"The real cost of the introduction of AI technology is not just the money some people pay for it. It is the social, political, and economic costs—to everybody in society that need to be considered."

4. EST-IL ACCEPTABLE DE CONFIER UNE PERSONNE VULNÉRABLE AUX BONS SOINS D'UNE IA ? (PAR EXEMPLE, AVEC UN « ROBOT-NANNY »)

Les répondants sont très partagés sur cette question : « pour divertir, mais pas pour soigner », « pas sûr ». On semble craindre une disparition de l'humain dans le soin. On pointe l'importance de « la chaleur humaine » en particulier pour les personnes vulnérables. Reste que c'est mieux que rien : « Oui, s'il n'y a pas d'autres choix ». Il faut aussi voir que cela pourrait donner un accès à de meilleurs soins, en particulier lorsque les ressources humaines manquent. Plusieurs notent cependant le risque de se déresponsabiliser de nos devoirs envers ces personnes en les confiant à une IA. Le sujet est sensible et de tels robots devraient certainement être encadrés et supervisés.

EXTRAITS CHOISIS

« Pas totalement. Le robot-nanny devrait toujours être là comme complément au personnel humain. »

« Oui, si c'est possible de bien programmer l'IA pour qu'elle n'outrepasse pas certaines compétences plus sensibles. »

« À la personne vulnérable de décider. »

"No. The result could be disastrous as it has not been studied for decades to determine the social, psychological, mental & physical effects it would have on our children. It could also possibly make our children emotionally unable to connect & bond with their parents, siblings & other humans."

"Of course... consider how television is sometimes referred to as a babysitter."

5. UN AGENT ARTIFICIEL COMME TAY, LE CHATBOT « RACISTE » DE MICROSOFT, PEUT-IL ÊTRE MORALEMENT BLÂMABLE ET RESPONSABLE ?

Une question qui suscite plutôt des réponses négatives. On refuse de qualifier le *chatbot* comme raciste « puisqu'il ne comprend rien » pour faire porter la responsabilité sur ses concepteurs (Microsoft). Il n'empêche que « les conséquences de ses déclarations » pourraient avoir des effets bien réels. La plupart des répondants sont donc d'accord pour y voir quelque chose d'inacceptable. Un répondant explore l'angle juridique en envisageant de mettre les IA sous tutelle légale (comme des enfants ou des animaux) tandis qu'un autre les envisage comme de simples objets dont la responsabilité incombe au propriétaire.

EXTRAITS CHOISIS

« Non, je crois que nous devrions considérer les produits de l'intelligence artificielle comme si c'était des enfants. Il serait pertinent de leur donner un titre de personne n'ayant pas la personnalité juridique autonome complète. Comme cela, chaque produit intelligent artificiellement aurait un humain qui serait tuteur responsable de ses actes. »

« Il ne s'agit en fin de compte que d'un programme. Et l'on sait jusqu'à quel point des programmes peuvent être bogués, déficients et mal faits. »

« Pas pour le moment, la responsabilité vient avec la conscience, si l'IA n'est pas consciente, elle ne peut être blâmable. »

« Il est de la responsabilité du programmeur de s'assurer que son robot ne soit pas raciste et d'apporter tout changement requis dans les plus brefs délais. »

"We should accept that machine learning algorithms are non-deterministic and empower operators to explore their utility while being responsible operators."

"The responsibility (until proven that the being is actually sentient, if that's even possible) should be taken by: People who gave permission to deploy them > People who tested them > People who developed them. In that order."

"Humans are not good examples for AI agents. AI agents will be more efficiently learning from other AI agents than from human activities."

"No. I think it is always people who must be held responsible. I am against giving machines any kind of legal status similar to people. You cannot punish or hold responsible a machine. So, people must always be responsible."

3. SYNTHÈSE DES MÉMOIRES REÇUS

Une quinzaine de documents ont été reçus à la suite de l'appel lancé via le site web de la Déclaration de Montréal en novembre 2017 (avec une date limite de réception fin avril 2018). Il s'agissait de contribuer au contenu de la Déclaration, soit en discutant les 7 principes de la version préliminaire, soit en suggérant des recommandations concrètes. Ces documents vont du mémoire de synthèse de discussions collectives au texte d'opinion individuel. Ils sont écrits en français et en anglais ; on peut les lire sur le site de la Déclaration (cette synthèse ne rend évidemment pas compte de toute la richesse des mémoires reçus).

Les abréviations suivantes seront utilisées pour désigner les documents des organismes ou personnes suivantes :

AQT
pour l'Association québécoise des technologies

CAIQ
pour la Commission d'accès à l'information du Québec

MAIEM
pour le groupe Montréal AI Ethics Meetup

OIQ
pour l'Ordre des ingénieurs du Québec

SRAD
pour la soirée de réflexion autour de la Déclaration qui s'est tenue à l'UQAM

Hernandez
pour Annick, Guillaume et Raphaël Hernandez

McNally
pour John McNally

Musseau
pour Pierre Musseau-Milesi

Parent pour Lise Parent

Quintal et al.
pour Ariane Quintal, Matthew Sample et Eric Racine

Ravet
pour Jean-Claude Ravet

Robert
pour Bruno Robert

Wark
pour Grant Wark

VIE PRIVÉE

PRINCIPE PROPOSÉ :

« Le développement de l'IA devrait offrir des garanties sur le respect de la vie privée et permettre aux personnes qui l'utilisent d'accéder à leurs données personnelles ainsi qu'aux types d'informations que mobilise un algorithme ».

REMARQUES GÉNÉRALES

Le principe relatif à la vie privée est probablement celui qui a été le plus longuement commenté dans les mémoires reçus. La Commission de l'accès à l'information du Québec (CAIQ) en particulier, mais aussi le groupe Montréal AI Ethics Meetup (MAIEM), le rapport de la soirée de réflexion autour de la Déclaration qui s'est tenue à l'UQAM (SRAD), l'Ordre des ingénieurs du Québec (OIQ), Lise Parent (Parent), Annick, Guillaume et Raphaël Hernandez (Hernandez), Grant Wark (Wark), Quintal, Sample et Racine (Quintal et al.) ont proposé des recommandations explicitement liées à la vie privée.

Comme le remarque CAIQ, au Québec, la protection des renseignements personnels possède déjà des principes bien établis (RLRQ, A-2.1 ; la Loi sur l'accès, ainsi que RLRQ, P-39.1 ; la Loi sur le privé) que le développement de l'IA devra respecter : par exemple, les organismes qui collectent des données doivent déterminer par avance la finalité de cette collecte et en tenir informées les personnes concernées. On peut encore nommer les principes de nécessité, de consentement, de confidentialité, de destruction, de transparence, d'accès et de responsabilité (voir l'annexe de CAIQ).

Face à de nouvelles pratiques, on peut envisager au moins deux types de régulations : l'une coercitive, qui met l'accent sur les pénalités en cas de non-respect du cadre légal, et l'autre, préventive, qui vise à accompagner les changements de façon

plus souple. Dans le contexte québécois, la CAIQ suggère le second type d'approche et insiste sur l'évaluation des risques en amont, sur l'utilisation de paramétrages qui soient, par défaut, les plus stricts possible, l'utilisation de technologie pour améliorer la confidentialité, la désignation dans chaque organisation d'un responsable à la protection des renseignements personnels qui soit imputable et la « transparence, au profit du citoyen ». On peut toutefois se demander si le rapport de force avec les grandes multinationales du numérique ne requerra pas aussi des mesures musclées de type plus coercitif que préventif.

Cette position peut faire écho à celle de l'OIQ (et de Parent) qui promeut la protection de la vie privée dès la conception (*privacy by design*) et propose de s'inspirer des bonnes pratiques existantes comme celles du règlement sur la protection des données (RGPD) récemment entré en vigueur en Europe.

Ce souci pour le respect de la vie privée s'accompagne souvent de celui pour la transparence. Le groupe MAIEM propose ainsi de bonifier le principe de vie privée en précisant que la transparence est essentielle — une analyse qui est aussi faite par CAIQ et SRAD. On retrouve donc ici la relation étroite entre les enjeux de protection de certaines informations (les renseignements personnels) et la possibilité de savoir qui détient quoi (l'accès à l'information), deux éléments qui mériteraient sans doute d'être plus explicites dans la Déclaration. On peut noter en passant que ces éléments peuvent entrer en tension : lorsque la transparence s'applique à des renseignements personnels qu'on voudrait garder confidentiels. Des arbitrages entre ces deux notions peuvent être nécessaires.

Il ne va d'ailleurs pas de soi que ces arbitrages soient consensuels, car comme le souligne MAIEM, les préférences en matière de vie privée peuvent « varier considérablement selon les cultures, les générations et les individus ». On constate en tout cas un consensus sur l'idée qu'il faut « préserver le contrôle du citoyen sur ses renseignements personnels et la gestion de son consentement » (CAIQ, SRAD). Quintal et al. s'inquiètent d'ailleurs de ce que la formulation initiale du principe de vie privée suggère que les données soient partagées par

défaut (le principe insiste sur la possibilité de savoir ce que deviennent les données personnelles, sans contester que ces données soient recueillies dans un premier temps). « La Déclaration, écrivent-ils, devrait inclure des garanties élevées pour la confidentialité des données des utilisateurs » [The Declaration should include improved safeguards for privacy of user data.]

SRAD note enfin que les techniques d'anonymisation des données ne sont pas encore matures pour être utilisables sans risque. SRAD remarque également le lien entre les enjeux de protection des données et les risques de discriminations algorithmiques. Mais cela ne signifie pas que des données protégées (par exemple le genre ou la race) ne devraient pas être recueillies dans la mesure où combattre une discrimination suppose habituellement d'avoir accès à ces informations.

PROPOSITIONS DE RECOMMANDATIONS

Le principe de vie privée, intégrant le souci de transparence, donne lieu à des recommandations plus spécifiques :

- > Les gens doivent être informés, autorisés et capables de vérifier leurs renseignements personnels et la manière dont ils sont utilisés à tout moment. [People need to be informed about, and allowed and able to check, their personal data and its uses at any time] (MAIEM).
- > Il faut introduire une culture du « data privacy by default » comme c'est le cas en neuro-éthique, c'est-à-dire que, par défaut, les données personnelles ne devraient pas être partagées (Quintal et al.).
- > Le « fardeau du consentement », c'est-à-dire le souci de s'assurer qu'il existe un véritable consentement libre et éclairé, devrait incomber aux entreprises/organisations et non aux citoyens, de même que celui de corriger des informations inexactes (CAIQ).
- > Les gens doivent pouvoir comprendre comment leurs renseignements personnels seront utilisés (MAIEM, CAIQ, Hernandez, Parent).

- > Les gens doivent pouvoir retirer leur consentement à l'utilisation de leurs renseignements personnels (MAIEM).
- > Il faut rendre publics et ouverts les codes informatiques pour l'interprétation des résultats et les méthodes d'entraînement des algorithmes. (OIQ)
- > Il faut sensibiliser les gens aux enjeux de protection de la vie privée (CAIQ).
- > Les gens devraient pouvoir connaître en tout temps la valeur monétaire de leurs renseignements personnels (Hernandez).

Enfin, une proposition originale et détaillée de Wark vient d'une certaine façon répondre à une interrogation de Hernandez : Comment créer un espace numérique privé ? Il s'agit en quelque sorte d'utiliser l'IA pour se protéger de l'IA.

- > En effet, Wark suggère d'utiliser la technologie des « *smart contracts* » pour protéger les renseignements personnels et faciliter les échanges commerciaux et les interactions sociales. Il s'agit de développer un profil personnel sécurisé et une « IA loyale » qui feraient office de gestionnaire des données personnelles, répondant ainsi à nombre de défis identifiés précédemment. "For example, a loyal AI-agent must not adulterate its loyalty to its owner through overt or covert association with a business, such as an online store." Pour en savoir davantage, nous renvoyons au mémoire de Wark qui présente de façon détaillée à quoi une telle IA loyale pourrait ressembler.

Plusieurs mémoires évoquent enfin la mise en œuvre de ces recommandations. Du point de vue des politiques publiques, on peut envisager au moins trois manières de traduire ce souci pour le respect de la vie privée et de la transparence : par la réglementation, par l'autoréglementation ou par des incitatifs.

Aussi bien CAIQ que MAIEM sont d'avis que l'autoréglementation ne peut être suffisante. Il importe plutôt de moderniser la réglementation existante. L'un et l'autre (ainsi que Parent) insistent aussi sur l'importance de faire des audits des

entreprises et organisations. Cette modernisation peut prendre différentes directions : l'OIQ soutient « des mécanismes règlementaires souples », ce qui résonne avec l'approche préventive défendue par CAIQ.

On peut enfin envisager des incitatifs économiques pour favoriser les entreprises qui développent des technologies protectrices de la vie privée ainsi qu'une valorisation de celles qui font des efforts, par exemple au moyen de labels ou de certifications — une idée qui semble aussi avoir les faveurs de l'Association québécoise des technologies (AQT).

JUSTICE

PRINCIPE PROPOSÉ :

« Le développement de l'IA devrait promouvoir la justice et viser à éliminer les discriminations, notamment celles liées au genre, à l'âge, aux capacités mentales et physiques, à l'orientation sexuelle, aux origines ethniques et sociales et aux croyances religieuses »

REMARQUES GÉNÉRALES

Comme le principe de vie privée, celui de justice a aussi été discuté ou évoqué dans de nombreux mémoires : MAIEM, SRAD, OIQ, Hernandez, Parent, McNally, Ravet.

SRAD propose de distinguer les différents sens de la justice (avec Aristote) : la justice commutative qui règle les échanges entre personnes considérées comme égales et la justice distributive qui s'attache au mérite. Qui dans la société mérite quoi ? C'est surtout cette seconde acception qui semble mobilisée dans les mémoires reçus. Et elle soulève de nombreuses questions.

Peut-on identifier un principe de justice universel pour régir le développement de l'IA ? Ne faudrait-il pas se contenter de principes valables uniquement pour une communauté donnée ? Cette question délicate est au cœur de nombreux débats en philosophie politique.

MAIEM penche pour une approche non universaliste, ou du moins, qui essaye de tenir compte des variations culturelles et historiques de la notion de la justice : « Le développement et l'utilisation de solutions fondées sur l'intelligence artificielle devraient promouvoir la justice et l'agentivité humaine, définies de manière transparente par l'organisation qui définit le bien-être de la communauté cible (par exemple le gouvernement démocratiquement élu), de concert avec la communauté cible. Il devrait chercher à éliminer les inégalités et la discrimination au sein de cette communauté. »

[The development and utilization of AI-enabled solutions should promote justice and human agency as transparently defined by the target community's welfare-defining organization (e.g. democratically elected government), in concert with the target community. It should seek to eliminate inequality and discrimination within that community.]

On peut opposer à cette reformulation l'approche plus universaliste de Ravet qui trouve un principe universel dans l'idée (kantienne) de dignité humaine et dans celle de vie : « Les innovations en IA doivent se fonder sur le principe de non-instrumentalisation de l'humain et veiller à ne pas écraser la vie. » C'est aussi l'approche de SRAD qui, en plus de la notion d'égale dignité humaine, introduit l'idée de justice sociale : « Le développement de l'IA devrait promouvoir la justice sociale et respecter l'égale dignité humaine, notamment en visant à éliminer toute forme de discrimination, incluant celle liée au genre, à l'âge, aux origines ethniques, aux statuts sociaux, etc. ».

Une manière d'articuler justice sociale et justice comme non-discrimination pourrait être de voir la première comme corrigeant des inégalités (socio-économiques) tandis que la seconde chercherait davantage à éviter que des inégalités n'apparaissent et à garantir l'égalité des chances. On peut

d'ailleurs aussi envisager la justice sociale dans une perspective plus contextuelle, comme le fait MAIEM qui insiste sur la nécessité de considérer différentes perspectives sur la justice, en particulier celles des communautés marginalisées.

PROPOSITIONS DE RECOMMANDATIONS :

La question des biais (déjà évoquée dans la section précédente) et de l'opacité des algorithmes (le « problème de la boîte noire ») a aussi retenu l'attention. Ce n'est pas forcément étonnant dans la mesure où cet enjeu a bénéficié d'une importante couverture médiatique. Parent note par exemple que « les systèmes de prise de décisions assistée, voire automatisée, en médecine, finance, défense ou justice, donneront des résultats biaisés si leurs intrants sont biaisés. » De même, l'OIQ insiste sur la nécessité d'instaurer des « mécanismes de contrôle et de protection » pour corriger les biais.

D'autres recommandations peuvent être mentionnées :

- > Il faut former les étudiants et praticiens en IA au droit et à l'éthique. (Parent, OIQ)
- > Il faut promouvoir l'emploi diversifié et féminin dans le développement de systèmes d'IA. (OIQ)
- > Il faut assurer un traitement rapide et transparent des réclamations des utilisateurs/citoyens qui auraient été affectés par les effets défavorables d'un système d'IA (OIQ).

Plusieurs mémoires appellent de leurs vœux la création d'un **organisme indépendant de supervision** (Parent, McNally, Hernandez, OIQ, AQT). Son rôle ne se cantonnerait pas à l'application du principe de justice, mais comme il est souvent mentionné à propos des enjeux de discriminations, on peut saisir l'occasion pour l'évoquer ici.

Sa présentation varie bien sûr selon les mémoires. L'OIQ parle d'un observatoire de l'IA, Hernandez évoque « un organisme régulateur dont la tâche serait d'assurer au citoyen une bonne compréhension des décisions prises par les IA » ; quant à l'AQT, elle préconise « la mise en place d'un comité des sages multisectoriel qui aura pour

mission d'engager une démarche réflexive sur les opportunités et les défis de l'industrie québécoise des technologies sur la question de l'éthique en intelligence artificielle. » On peut aussi penser comme le suggère McNally à un organisme de surveillance qui collaborerait étroitement avec le gouvernement et aurait pour mandat d'anticiper les problèmes que l'IA posera à la société de demain.

RESPONSABILITÉ

PRINCIPE PROPOSÉ :

« Les différents acteurs du développement de l'IA devraient assumer leur responsabilité en œuvrant contre les risques de ces innovations technologiques ».

REMARQUES GÉNÉRALES

Le principe de responsabilité est moins mentionné que les précédents dans les mémoires reçus, mais on peut dire que son ombre plane dès que la question de la relation entre les humains et l'IA est soulevée. Qui sera responsable de l'IA, en particulier de ses effets néfastes ? Comme le remarque SRAD, le développement de l'IA pourrait aller jusqu'à l'utilisation de robots tueurs. Cette possibilité soulève à son tour une inquiétude largement partagée : que les humains se déchargent de leurs responsabilités sur le dos de l'IA. On retrouve ici le thème de l'IA comme outil : celle-ci devrait être vue comme une extension de l'intentionnalité humaine, mais non comme une intentionnalité autonome (MAIEM).

Parmi les personnes et groupes responsables, on peut citer les chercheurs qui, détenant la connaissance, doivent soulever le débat (SRAD). Il faut y ajouter la responsabilité de ceux qui commanditent les chercheurs, comme les universités, les militaires ou les industries. Être

responsable signifie notamment mettre en place les savoirs et les outils pour « comprendre le fonctionnement de l'IA et anticiper ses réactions » (MAIEM).

Dans un mémoire qui offre davantage une perspective générale sur la conception de l'IA qui prévaut aujourd'hui qu'une bonification de tel ou tel principe de la Déclaration, Jean-Claude Ravet, le rédacteur en chef de la revue Relations, entend nous mettre en garde contre l'instrumentalisation de l'humain à l'ère de l'IA et estime que le développement de l'IA engage notre responsabilité collective et qu'il est nécessaire d'en avoir une vue d'ensemble, historique et idéologique. Ainsi, c'est le motif même de l'IA comme outil qu'il convient d'interroger, puisque « la frontière entre l'usage de la technique et la technique elle-même se brouille plus que jamais ». Surtout, note Ravet, il convient de ne pas être dupe de l'idéologie qui accompagne le développement de l'IA et sert les intérêts de puissances multinationales. Pour Ravet, cette idéologie qui entend se faire passer pour un discours scientifique plutôt que pour un projet de société assumé, se caractérise par « une vision extrêmement réductrice de l'humain et de la vie ». (Hernandez s'interroge aussi sur la spécificité de l'humain).

Le mouvement transhumaniste ou le livre *Homo Deus* de Yuval Noah Harari seraient des représentants de cette idéologie réductionniste que condamne Ravet : « Sous prétexte d'augmenter l'humain, on ne doit pas le diminuer et en faire un moyen en vue d'une fin. Le seul critère de rentabilité ne suffit pas. Ni le respect du choix individuel. Car les enjeux touchent au vivant et à l'humanité en tant que tels. » Il importe de jeter un regard critique sur ce qui apparaît bien souvent comme une évidence, à savoir que le progrès de l'humanité passe par l'IA et qu'il y a quelque chose d'inévitable dans le déploiement de ces technologies. Autrement dit, c'est collectivement que nous sommes responsables et c'est bien l'humain qui doit toujours garder le dernier mot, « en tant qu'être de parole, de sentiments, de sensations, et conscient de sa fragile humanité et des liens qui l'unissent aux autres, au vivant et à la Terre » (Ravet).

PROPOSITIONS DE RECOMMANDATIONS

- > Les décisions de justice faites avec l'assistance d'une IA devront ultimement être imputables à des êtres humains (SRAD, Parent, Ravet).
- > Dans le cas des ingénieurs, il faut préserver l'imputabilité des professionnels dans l'exercice de leur profession (OIQ).
- > Du point de vue de la responsabilité juridique, il faut anticiper d'éventuels litiges autour des systèmes d'IA avec les juridictions non canadiennes (ex. composantes conçues ou fabriquées ailleurs que là où le système a été utilisé) (OIQ).
- > Pour éviter d'attribuer une responsabilité induite aux IA, elles ne devraient pas avoir l'apparence trompeuse d'un patient moral (c'est-à-dire un individu pouvant subir un tort) qui mérite notre empathie (MAIEM).
- > Dans la formulation du principe, l'intention de « lutter contre les risques » n'est pas assez exigeante : les responsables devraient assumer les résultats du développement de l'IA (MAIEM).

BIEN-ÊTRE

PRINCIPE PROPOSÉ :

« Le développement de l'IA devrait ultimement viser le bien-être de tous les êtres sentients. »

REMARQUES GÉNÉRALES

Comme la responsabilité, le principe du bien-être est souvent présent de manière implicite, en particulier lorsqu'il est question de santé, de sécurité ou même de juste répartition des bénéfices de l'IA. En fait, selon certaines approches en philosophie morale, c'est même un principe qui pourrait servir de critère général de prise de décision : lorsqu'on

a le choix, on devrait agir de manière à produire le plus de bien-être possible. Évidemment, comme le remarque notamment MAIEM, d'autres valeurs peuvent entrer en conflit avec le bien-être, en particulier l'autonomie. Par exemple, on peut décrire les situations où le paternalisme semble justifié comme des situations où on restreint l'autonomie d'un patient moral au nom de son bien-être. Dès lors, il n'est pas étonnant que la possibilité de tels conflits de valeurs — souvent discutés par les philosophes dans les dilemmes moraux — soit envisagée dans les mémoires reçus. Pour autant, il n'en demeure pas moins vrai qu'un principe portant sur le bien-être se doit d'être simple, facile à comprendre et laisser une certaine marge pour des interprétations futures (MAIEM).

De son côté, l'OIQ remarque que le principe du bien-être s'accorde bien avec une des principales dispositions du code de déontologie de l'ingénieur (article 2.02) qui stipule que « l'ingénieur doit respecter ses obligations envers l'homme et tenir compte des conséquences de l'exécution de ses travaux sur l'environnement et sur la vie, la santé et la propriété de toute personne. » C'est pourquoi promouvoir le bien-être suppose d'évaluer le plus finement possible les risques liés au déploiement et au fonctionnement des applications de l'IA, en gardant en tête que « le risque zéro n'existe pas » (OIQ).

Il faut aussi noter que le caractère très inclusif de ce principe, qui vise non seulement le bien-être des humains, mais aussi de l'ensemble des êtres sentients n'a pas été remis en question. C'est peut-être le signe d'une évolution des mentalités quant à notre rapport aux animaux non humains (sentients). MAIEM et SRAD reprennent à leur compte cette extension du domaine de la moralité aux êtres sentients tandis que Parent évoque les interférences de l'IA avec la vie animale. On peut aussi constater que certains mémoires (Ravet, MAIEM, Parent) semblent vouloir tenir compte du critère de la vie pour étendre le cercle de la moralité à des entités non sentientes (comme des végétaux ou des écosystèmes). Ces intuitions qu'on pourrait qualifier de biocentristes ne sont toutefois pas assez développées pour qu'on puisse dire qu'elles reflètent une position morale (assez radicale) assumée : il peut s'agir d'un souci pour l'environnement d'ordre anthropocentriste.

L'idée semble aussi partagée que la capacité pour une IA d'être sentiente (ou sensible) serait un bon critère pour lui conférer des droits ou, à tout le moins, de la considération morale puisque si un robot pouvait souffrir, par exemple, il aurait alors un intérêt légitime à ce qu'on le protège. Ce point demeure toutefois très spéculatif dans la mesure où les systèmes d'IA sont actuellement très loin d'éprouver des sensations ou des émotions.

Enfin, dans un texte assez spéculatif et programmatique, Museau cherche à articuler la notion de minimalisme morale développée par le philosophe Ruwen Ogien et le principe du bien-être proposé par la Déclaration. Il en ressort que le développement de l'IA devrait viser à ne pas faire de tort à autrui et non pas à s'améliorer soi-même — l'amélioration de soi étant, selon Museau, à la fois de l'ordre du maximalisme moral et du transhumanisme.

PROPOSITIONS DE RECOMMANDATIONS

- > **Formuler le principe en termes de réduction des souffrances plutôt que de promotion du bien-être (ce qui correspond à ce qu'on nomme parfois l'utilitarisme négatif) (MAIEM).**
- > **Par souci de sécurité, il faut prévoir des dispositifs de débrayage/blocage dès la conception des systèmes d'IA afin d'en conserver le contrôle en cas de défaillance (OIQ).**

AUTONOMIE

PRINCIPE PROPOSÉ :

« Le développement de l'IA devrait favoriser l'autonomie de tous les êtres humains et contrôler, de manière responsable, celle des systèmes informatiques ».

REMARQUES GÉNÉRALES

Concernant l'autonomie, on peut dire qu'il y a un consensus quant à promouvoir l'autonomie des humains. Cette idée se traduit notamment par le thème déjà évoqué de l'IA *au service* des humains. L'OIQ note ainsi que « les robots et les systèmes d'IA doivent être vus comme des outils d'assistance ou d'aide à la décision et non comme une substitution au jugement humain ». De son côté, Ravet insiste pour ne pas réduire l'humain à une machine ni en faire un moyen en vue d'une fin, tandis qu'Hernandez se demande si l'IA ne va pas un jour remplacer les humains de sorte qu'ils deviennent obsolètes.

Il n'en demeure pas moins vrai que la notion d'autonomie renvoie à de multiples acceptions. SRAD propose ainsi une grille d'analyse détaillée des types d'autonomie (« condition d'une entité qui choisit elle-même les lois auxquelles elle se soumet ») divisée en autonomie morale, politique et fonctionnelle (non-dépendance). Ces trois types d'autonomie peuvent être croisés avec trois types de situations : l'autonomie d'un humain aidé par une IA (par exemple une personne en situation de handicap), l'autonomie d'un humain dans un environnement peuplé d'IA et, enfin, l'autonomie d'une IA dans un environnement humain. SRAD propose dès lors une reformulation qui tienne davantage compte de ces diverses acceptions : « Les systèmes IA ne doivent pas nuire à l'autonomie (morale, fonctionnelle, politique) des êtres humains, mais devraient chercher à y contribuer. Les systèmes IA ne doivent pas être rendus entièrement autonomes des êtres humains, mais doivent demeurer sous leur contrôle (moral, fonctionnel, politique) ». Pour autant,

on n'en déduira pas trop vite que l'autonomie devrait systématiquement prévaloir sur les autres valeurs comme le bien-être, la justice ou la connaissance. Chaque cas est à examiner en contexte. Et comme le rappelle MAIEM, le consentement des gens demeure une bonne manière de garantir leur autonomie.

Si la valeur de l'autonomie humaine fait consensus, la situation est plus délicate pour ce qui a trait à « l'autonomie des systèmes d'IA » dans la mesure où leur mise sous tutelle pourrait être contestée. Ainsi, citant un article sur l'évolution digitale et la vie artificielle, MAIEM rappelle qu'on pourrait imaginer des situations dans lesquelles l'autonomie et la créativité des systèmes d'IA seraient profitables au bien-être général. Il n'empêche, poursuit MAIEM, que l'autonomie d'un système d'IA ne devrait pas être recherchée pour elle-même si cela entre en conflit avec le bien-être d'un être sentient. Pour pertinentes qu'elles soient, ces remarques sont assez isolées parmi les mémoires : ceux-ci donnent plutôt le sentiment qu'il faut étroitement surveiller les systèmes d'IA au risque d'en perdre le contrôle. Il semble toutefois envisageable de concilier ces considérations apparemment divergentes : on pourrait garder le contrôle à un certain niveau sur un système d'IA tout en autorisant — à un plus bas niveau et dans un cadre déterminé — que l'IA expérimente certaines solutions à des problèmes de façon libre et créative.

PROPOSITIONS DE RECOMMANDATIONS

- > La notion d'autonomie a davantage suscité des réflexions philosophiques que des recommandations concrètes, même si certaines recommandations des autres sections ne sont pas sans rapport (par exemple celles sur le consentement dans la section vie privée).

CONNAISSANCE

PRINCIPE PROPOSÉ :

« Le développement de l'IA devrait promouvoir la pensée critique et nous prémunir contre la propagande et la manipulation ».

REMARQUES GÉNÉRALES

Les liens entre l'IA et la connaissance sont multiples. Tout d'abord, et dans la perspective des sciences cognitives, l'intelligence artificielle peut nous aider à comprendre l'intelligence naturelle, l'une et l'autre pouvant se définir comme ce qui guide une capacité d'action (SRAD). Dès lors, on peut se demander pourquoi l'intelligence naturelle devrait prévaloir sur l'intelligence artificielle puisqu'à un certain niveau d'analyse, les humains ou les animaux, tout comme les machines, sont des systèmes causaux.

Dans plusieurs mémoires, le principe de connaissance est l'occasion de discuter des enjeux de propagande et des fausses informations (*fake news*). De ce point de vue, le sujet concerne tout autant la démocratie que la connaissance. Ainsi, on peut se demander comment une IA ou ceux qui la produisent et la commercialisent pourraient être en position de décider ce qui relève de la propagande ou de la manipulation. Il semble illégitime, voire dangereux de leur confier une telle responsabilité. C'est pourquoi MAIEM propose une reformulation du principe qui mette plutôt l'accent sur la transparence : « Le développement de l'IA ne devrait pas nuire à la pensée critique. Il devrait également être transparent et ouvert afin de permettre la participation du public au développement de l'IA, un contrôle public de l'IA et l'éducation à l'IA. »
[*The development of AI should not hamper critical thinking. It must also proceed in a transparent and open manner, to enable public participation in its development, scrutiny, and education*].

Parmi les autres thèmes liés à la connaissance, on trouve l'accès public aux résultats des recherches en IA, la pensée critique (MAIEM nous met en garde contre les chambres d'écho), l'éducation à l'IA et l'opacité des algorithmes déjà mentionnés dans la section justice. Sur ce dernier point, SRAD en appelle à des efforts pour améliorer la transparence des données et des algorithmes, mais aussi pour publier les codes sources derrière les IA.

PROPOSITIONS DE RECOMMANDATIONS

- > Des mesures devraient être mises en place afin de promouvoir l'accès public aux résultats des recherches universitaires. (MAIEM)
- > Il faut encourager la compétition et la diversité dans les applications de l'IA afin que cela bénéficie à toute la société. (MAIEM)
- > Il faut repenser le *business* modèle des médias sociaux des autres sites de nouvelles. (MAIEM)
- > Tous les étudiants et praticiens en IA devraient recevoir une formation avancée en éthique. (Parent)

DÉMOCRATIE

PRINCIPE PROPOSÉ :

« Le développement de l'IA devrait favoriser la participation éclairée à la vie publique, la coopération et le débat démocratique »

REMARQUES GÉNÉRALES

Concernant la démocratie, plusieurs mémoires (Robert, Parent, OIQ, AQT, SRAD) saluent l'initiative de la Déclaration et la possibilité qu'elle leur offre

de faire entendre leur point de vue. MAIEM y voit une « contribution importante » aux discussions internationales sur le sujet.

D'autres sont plus critiques. Ainsi, Quintal et al. contestent le processus même de production de la Déclaration de Montréal. S'ils sont favorables aux efforts de consultation publique, ils se demandent toutefois s'il ne s'agit pas là d'une entreprise de légitimation visant à entériner un document existant. Plus précisément, ils craignent que la version préliminaire de la Déclaration (soit les 7 principes proposés qui servent d'articulation à cette synthèse) n'ait fortement orienté les débats citoyens : « Le public aurait dû être impliqué dans la délibération sur le contenu de la Déclaration dès le tout début » [the public should have been meaningfully engaged in deliberating on the contents of the Declaration from the very beginning]. Pour Quintal et al., cela risque de compromettre la légitimité finale de la Déclaration.

Ces inquiétudes signifient évidemment un appel à plus de démocratie (et de transparence et d'esprit critique) dans le développement de l'IA, ce qui appuie en quelque sorte le principe lié à la démocratie. En outre, précisent Quintal et al., la bonne volonté démocratique restera un vœu pieux si elle ne s'accompagne pas d'une régulation de l'industrie. On court aussi le risque que les compagnies utilisent les algorithmes pour confiner le débat à des enjeux qui leur paraissent acceptables (ce qu'on pourrait avec SRAD qualifier d'enjeu épistémique ayant des effets néfastes pour la démocratie). Un argument similaire est proposé par MAIEM qui note que puisqu'il est peu probable que les compagnies partagent leurs algorithmes afin de protéger leur propriété intellectuelle, la régulation externe semble être la meilleure solution.

Sur le principe lui-même, MAIEM trouve sa formulation un peu vague et regrette qu'il se concentre sur la démocratie alors que tous les humains ne vivent pas dans de tels régimes. MAIEM suggère dès lors de lui substituer un « principe de participation publique » qui se présenterait ainsi : « Le développement de l'IA devrait s'accompagner d'une information claire et précise afin de permettre un débat éclairé sur l'IA et ses applications, tout en encourageant l'ouverture et la transparence dans la recherche. » [The development of AI should promote

the dissemination of clear and accurate information to the public to enable open and educated debate about AI and its applications, and encourage open and transparent research collaboration.]

SRAD, enfin, rappelle que les grandes entreprises de technologie (comme les GAFA) possèdent aujourd'hui un pouvoir considérable, tant économique que politique — notamment parce qu'elles ont un accès direct à énormément de données personnelles. On peut y voir une menace sérieuse pour la démocratie comme le suggère l'affaire Cambridge Analytica qui a éclaté depuis lors. Par ailleurs, dans la mesure où la démocratie requiert une certaine égalité socioéconomique — au risque de dégénérer en oligarchie — il faut prendre garde à l'accroissement des inégalités que devrait mécaniquement entraîner le développement de l'IA. En effet, explique SRAD, l'automatisation d'une tâche par une IA revient à déplacer la richesse du revenu vers le capital (et donc à la concentrer entre les mains des actionnaires plutôt que des salariés remplacés par l'IA). À moins de régulations ou d'encadrement, l'IA risque donc d'amplifier la croissance des inégalités économiques qu'on constate depuis les années 1950.

PROPOSITIONS DE RECOMMANDATIONS :

- > Les chercheurs à l'avant-garde du secteur dans nos institutions publiques universitaires doivent garder leur indépendance face à l'entreprise privée. (Parent)
- > Il faut briser les grands monopoles dans l'industrie technologique. (SRAD)
- > Il faut sérieusement considérer la possibilité d'un revenu universel garanti financé par une taxe sur l'automatisation ou le capital. (SRAD)
- > Il faut encourager de nouvelles structures de propriété d'entreprises telles que des coopératives pour lutter contre la concentration du capital. (SRAD)

CRÉDITS DU RAPPORT FINAL

Le rapport de la Déclaration de Montréal IA responsable a été rédigé sous la direction de :

Marc-Antoine Dilhac, instigateur du projet et responsable du Comité d'élaboration de la Déclaration ; codirecteur scientifique de la coconstruction ; professeur au Département de philosophie de l'Université de Montréal ; chaire de recherche du Canada en Éthique publique et théorie politique ; directeur de l'axe Éthique et politique, Centre de recherche en éthique

Christophe Abrassart, codirecteur scientifique de la coconstruction, professeur à l'École de design et codirecteur du Lab Ville Prospective à la Faculté de l'aménagement de l'Université de Montréal, membre du Centre de recherche en éthique

Nathalie Voarino, coordonnatrice scientifique de l'équipe de la Déclaration, candidate au doctorat en bioéthique, Université de Montréal

Coordination

Anne-Marie Savoie, conseillère, vice-rectorat à la recherche, à la découverte, à la création et à l'innovation de l'Université de Montréal

Collaboration aux contenus

Camille Vézy, candidate au doctorat en communication, Université de Montréal

Révision et édition

Chantal Berthiaume, gestionnaire de contenu et rédactrice

Anne-Marie Savoie, conseillère, vice-rectorat à la recherche, à la découverte, à la création et à l'innovation de l'Université de Montréal

Joliane Grandmont-Benoit, coordonnatrice de projets, vice-rectorat aux affaires étudiantes et aux études, Université de Montréal

Traduction

Rachel Anne Normand et François Girard, Services linguistiques Révidaction

Graphisme

Stéphanie Hauschild, directrice artistique

La rédaction de ce rapport n'aurait pu être possible sans les réflexions des citoyens, des professionnels et des experts ayant participé aux ateliers.

NOS PARTENAIRES

Université 
de Montréal



CENTRE DE RECHERCHE EN ÉTHIQUE



ICRA
Programme
IA et
société



Québec 
Fonds de recherche – Nature et technologies
Fonds de recherche – Santé
Fonds de recherche – Société et culture



Centre
Culturel
Canadien
Paris

Canadian
Cultural
Centre
Paris





</>

declarationmontreal-iaresponsable.com