



Avril 2026

Co-concevoir une politique organisationnelle d'utilisation de l'IA centrée sur la justice

Notre quatrième séance de la Série d'apprentissage Équité et IA portait sur la façon dont les organisations peuvent prendre des décisions intentionnelles et alignées sur leurs valeurs concernant l'utilisation de l'IA. Animée par la Dre Nasreen Rajani (elle/elle) (<https://nasreenrajani.com/>), la séance a exploré comment dépasser les principes éthiques généraux de l'IA pour parvenir à des orientations pratiques au niveau organisationnel – notamment par l'élaboration de politiques d'utilisation de l'IA. Une perspective centrée sur la justice a guidé la discussion, combinant une courte présentation et un processus interactif de co-conception à l'aide d'un tableau Miro. Les participantes et participants ont partagé de véritables défis, réfléchi à leurs pratiques actuelles et contribué à définir ce que devraient inclure des politiques d'IA significatives et utilisables. L'accent était mis sur le passage de principes abstraits à des politiques pratiques permettant aux organisations de guider leurs décisions, de réduire les risques et de rester alignées sur leurs valeurs.

Pourquoi cette séance est importante

Les organisations sont poussées à adopter l'IA – souvent rapidement, et sans orientation claire. De nombreuses équipes utilisent déjà l'IA de manière informelle : rédaction de contenu, analyse de données, appui aux propositions. Parallèlement, il existe de réelles préoccupations concernant les biais, la confidentialité, l'utilisation des données et les préjudices sociaux et environnementaux graves. La question n'est plus simplement « Devons-nous utiliser l'IA ? ». Elle est désormais :

1. À quelles fins l'utilisons-nous ?
2. Dans quelle mesure s'aligne-t-elle – ou entre-t-elle en conflit – avec nos valeurs ?
3. Qui participe à ces décisions ?
4. À quels usages refusons-nous de l'employer ?

Ce sont des questions de responsabilité, de pouvoir et d'imputabilité – en particulier pour les organisations qui travaillent avec les communautés les plus touchées par les préjudices. Les politiques d'utilisation de l'IA constituent l'un des moyens les plus pratiques de combler ce fossé. Lorsqu'elles sont bien conçues, elles :

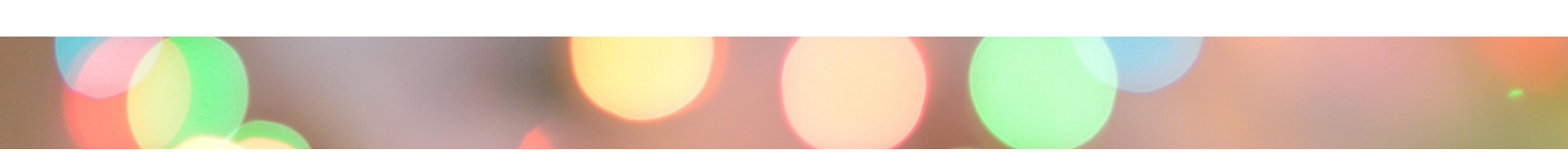
1. Fournissent des orientations claires au personnel
2. Créent de la cohérence entre les équipes
3. Établissent des limites face aux risques
4. Favorisent la transparence avec les partenaires et les communautés

Cette séance portait sur la façon de créer des politiques d'utilisation de l'IA qui soient utilisables, ancrées dans la réalité et alignées sur les valeurs organisationnelles.

Ce avec quoi les organisations peinent

La discussion de groupe a mis en lumière des défis concrets que les organisations rencontrent lorsqu'elles tentent de mettre en œuvre des politiques – pas seulement des politiques d'IA, mais des politiques en général. Ces constats pointent directement vers ce qui rend les politiques utilisables – ou non. Les participantes et participants ont soulevé plusieurs défis récurrents :

1. Les politiques deviennent rapidement obsolètes à mesure que les outils et les usages évoluent
2. Manque de clarté sur ce à quoi ressemble concrètement une « utilisation responsable de l'IA »
3. Les politiques ne sont pas intégrées dans les flux de travail quotidiens
4. Appropriation limitée dans les équipes : les politiques sont élaborées mais pas activement utilisées
5. Difficulté à intégrer le nouveau personnel et à maintenir l'alignement de toutes et tous
6. Les politiques ne sont pas communiquées de manière cohérente ni réévaluées régulièrement



Guide pratique : Comment élaborer une politique d'IA dans votre organisation

Au cours de la séance, nous avons exploré les étapes clés de l'élaboration d'une politique d'utilisation de l'IA – ce qu'elle doit inclure, comment l'aborder et ce qui la rend pratique et utilisable. Le guide ci-dessous est une synthèse de ces discussions, s'appuyant sur les réflexions du groupe, l'expérience de Nasreen en matière de conception de politiques d'IA centrées sur la justice, et les éléments clés de la Trousse d'outils de politique d'IA pour les OBNL canadiens : <https://mitchellconsultingsolutions.ca/resources/ai-policy-toolkit/#toolkit>

Étape 1 : Ancrer dans votre mission et vos objectifs (pas dans les outils)

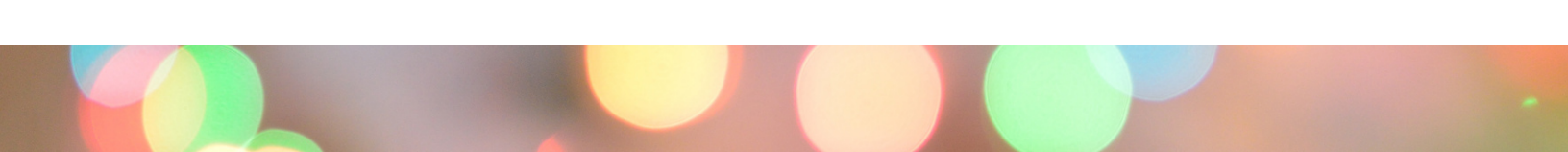
Avant de rédiger quoi que ce soit, commencez par clarifier pourquoi vous utilisez l'IA et en quoi cela est lié à votre mission, vos valeurs et vos responsabilités. Cela comprend l'identification des problèmes que vous cherchez à résoudre, la compréhension des domaines où l'IA peut apporter une valeur ajoutée, et là où elle pourrait introduire des risques ou des tensions – en particulier eu égard aux communautés que vous servez. Un point clé soulevé lors de la séance est que ce processus ne devrait pas se dérouler en vase clos. Les approches centrées sur la justice en matière de conception de politiques d'IA mettent l'accent sur la création d'espaces permettant aux communautés, à la société civile et aux groupes touchés de façonner les décisions en se basant sur leur vécu – et non sur la seule expertise technique. Par exemple, des initiatives comme la Consultation populaire sur l'IA (<https://www.peoplesaiconsultation.ca>) proposent des guides de facilitation structurés, des formats de discussion et des questions directrices pour soutenir un développement de politiques inclusif et participatif. Ancrer votre démarche de cette façon contribue à s'assurer que l'utilisation de l'IA est non seulement alignée sur votre mission, mais également imputable envers les personnes les plus touchées. Si l'IA ne soutient pas clairement votre travail – ou ne peut pas être utilisée d'une manière qui reflète vos valeurs – il n'est peut-être pas approprié d'en faire une priorité.

Étape 2 : Définir vos principes fondamentaux

Les politiques d'IA solides reposent sur un petit nombre de principes directeurs clairs qui reflètent les valeurs de votre organisation et servent de points d'ancrage pour la prise de décision. Ces principes constituent une référence lors de la navigation dans l'incertitude ou face à des arbitrages. Les principes courants incluent : l'utilisation centrée sur l'humain (l'IA soutient les personnes sans les remplacer), l'équité et la justice (identification et atténuation actives des biais), la transparence (ouverture quant au moment et à la manière dont l'IA est utilisée), la protection de la vie privée et des données, et l'imputabilité (veiller à ce que les personnes restent responsables des décisions et des résultats). Ce qui importe le plus, ce n'est pas d'adopter une longue liste, mais de choisir quelques principes qui ont du sens dans votre contexte et qui peuvent être appliqués en pratique. Vous pourrez vous inspirer de cadres existants – tels que les Principes de l'IA de l'OCDE ou les travaux de la Fondation Mozilla – comme point de départ, puis les adapter pour refléter votre mission et les réalités des communautés que vous servez. Garder cette section concise (généralement 4 à 6 principes) la rend plus facile à communiquer, à retenir et à appliquer dans les équipes.

Étape 3 : Cartographier les usages actuels de l'IA

Avant de définir des règles, il est important de comprendre comment l'IA est déjà utilisée dans votre organisation.



De nombreuses équipes expérimentent déjà – souvent de manière informelle – des outils de rédaction de contenu, de résumé d’informations, d’analyse de données ou d’appui aux propositions. Cartographier ces usages permet de passer des préoccupations abstraites aux décisions concrètes. Une approche pratique consiste à créer un espace partagé simple (comme un document ou un fil de discussion Slack) où le personnel peut noter comment il utilise l’IA, ce qui fonctionne bien et où il ressent de l’incertitude ou de l’inconfort. Ces exemples font souvent émerger des schémas importants : là où l’IA apporte une valeur ajoutée, là où les résultats nécessitent des corrections importantes, ou là où il peut y avoir des risques liés à l’utilisation des données, aux biais ou aux erreurs de représentation. Cette étape est également l’occasion d’identifier les lacunes dans la compréhension de l’équipe et de faire remonter les différences d’accès ou d’aisance avec les outils d’IA. Ancrer votre politique dans ces expériences réelles la rend beaucoup plus pertinente et garantit que les orientations que vous développez reflètent la façon dont les gens travaillent réellement – et non la façon dont vous supposez qu’ils travaillent.

Étape 4 : Établir des limites claires d’utilisation (et de non-utilisation)

Une fois que vous avez compris comment l’IA est utilisée, l’étape suivante consiste à définir des limites claires. C’est là que de nombreuses politiques échouent – en s’appuyant sur des formulations vagues comme « utiliser de manière responsable », ce qui laisse trop de place à l’interprétation. Les politiques sont plutôt les plus utiles lorsqu’elles définissent clairement ce qui est autorisé, ce qui nécessite de la prudence et ce qui n’est pas permis – tant pour les cas d’utilisation que pour les outils. Une approche pratique consiste à définir des cas d’utilisation spécifiques et à les catégoriser. Par exemple, utiliser l’IA pour réfléchir à des idées ou structurer du contenu peut être approprié, tandis que le résumé de documents internes peut nécessiter de la prudence et une révision. D’autres usages – tels que le téléversement de transcriptions d’entretiens, le traitement de données personnelles ou sensibles, ou la génération de contenu représentant des voix communautaires – peuvent être explicitement interdits. Ces distinctions aident le personnel à comprendre non seulement quoi faire, mais aussi quand et pourquoi certains usages sont plus risqués.

Il est également important de considérer les limites liées aux outils eux-mêmes. Les outils d’IA ne sont pas tous équivalents dans leur façon de traiter les données, de stocker les informations ou d’entraîner les modèles. Certaines organisations choisissent de limiter l’utilisation à un petit nombre d’outils approuvés, ou de restreindre des fonctionnalités spécifiques (comme le téléversement de fichiers ou l’utilisation d’outils qui conservent les données). Être explicite sur les outils pouvant être utilisés – et dans quelles conditions – réduit l’incertitude et aide à gérer les risques.

Enfin, des politiques solides définissent également des zones de non-utilisation. Dans certains cas, la décision la plus responsable est de ne pas utiliser l’IA du tout – en particulier dans des contextes impliquant des données sensibles, des informations qualitatives complexes, ou des situations où l’exactitude, le consentement et la représentation sont essentiels. Nommer ces limites clairement renforce l’idée que l’utilisation responsable de l’IA inclut de savoir quand ne pas l’utiliser. L’objectif n’est pas nécessairement de restreindre tout usage, mais de créer suffisamment de clarté pour que le personnel puisse prendre des décisions éclairées et cohérentes. Lorsque les limites sont spécifiques, ancrées dans des scénarios réels et clairement communiquées, elles réduisent les risques et favorisent une utilisation plus réfléchie et imputable de l’IA.

Étape 5 : Clarifier les rôles et les responsabilités

Des limites claires ne sont utiles que si l'on sait également qui est responsable de leur application. L'un des défis récurrents soulevés lors de la séance est que les politiques existent souvent sans propriété claire – ce qui entraîne un usage inégal, de l'incertitude et un suivi limité. Définir les rôles et les responsabilités permet de s'assurer que l'utilisation de l'IA est cohérente, imputable et soutenue à l'échelle de l'organisation. Au minimum, les politiques devraient énoncer les attentes envers tout le personnel utilisant des outils d'IA. Cela comprend généralement l'utilisation des outils conformément à la politique, la vérification de l'exactitude des résultats et l'exercice du jugement – en particulier dans les contextes à risque élevé. Il est également important de clarifier le rôle de la direction dans le soutien à la mise en œuvre, notamment pour l'intégration du nouveau personnel, le renforcement des attentes et la supervision de l'utilisation de l'IA au sein des équipes.

Certaines organisations identifient également un rôle spécifique – souvent désigné comme champion ou championne de la politique – pour aider à maintenir la politique active dans le temps. Cette personne (ou ce petit groupe) peut soutenir la formation, recueillir les commentaires du personnel, suivre les problèmes émergents et coordonner les mises à jour au fur et à mesure de l'évolution des outils et des pratiques. Dans les équipes plus restreintes, ces responsabilités peuvent être partagées, mais les rendre explicites aide à prévenir que la politique devienne inactive ou négligée. Clarifier les rôles renforce également un principe important : bien que l'IA puisse soutenir le travail, la responsabilité des décisions et des résultats demeure entre les mains des personnes. Rendre cela visible dans la politique contribue à bâtir une culture d'imputabilité, où le personnel se sent soutenu dans l'utilisation de l'IA – mais comprend également son rôle pour s'assurer qu'elle est utilisée de manière appropriée.

Étape 6 : Identifier et aborder les risques clés

Une politique d'IA solide rend les risques visibles et actionnables. Bien que de nombreux risques soient déjà compris en principe, ils demeurent souvent implicites ou inégalement traités dans la pratique. Les nommer clairement aide les équipes à reconnaître là où une attention particulière est nécessaire et favorise une prise de décision plus cohérente à l'échelle de l'organisation. Plusieurs domaines de risque clés ont émergé lors de la séance : les biais et la discrimination dans les résultats de l'IA, particulièrement dans des contextes variés ou avec des groupes sous-représentés; la protection des données et la sécurité, surtout lorsqu'il s'agit d'informations personnelles, sensibles ou non publiées; et le risque d'erreur de représentation – où le contenu généré par l'IA simplifie, déforme ou parle inexactement au nom des communautés. Les participantes et participants ont également noté le risque de dépendance excessive à l'IA, où les résultats sont acceptés sans examen suffisant ni esprit critique.

Aborder ces risques ne nécessite pas de systèmes excessivement complexes, mais exige de la clarté. En pratique, cela signifie lier les risques à des orientations spécifiques dans la politique. Par exemple, si les biais sont une préoccupation, la politique peut exiger une révision et une validation humaines des résultats. Si la protection des données est un risque, elle peut clairement interdire la saisie de certains types d'informations dans des outils d'IA. Lorsque l'erreur de représentation est une préoccupation, il peut être nécessaire de limiter ou d'éviter l'utilisation de l'IA dans des contextes où la voix, le vécu et la nuance sont essentiels. Il est important de noter que les risques ne sont pas les mêmes pour toutes les organisations ou tous les contextes. Les politiques devraient refléter les réalités spécifiques de votre travail, notamment les populations que vous servez et les types de données que vous traitez. Prendre le temps d'identifier et d'articuler ces risques contribue à s'assurer que l'utilisation de l'IA est non seulement efficace, mais aussi responsable et alignée sur vos valeurs.

Étape 7 : Établir une gouvernance et des processus décisionnels simples

Une fois les limites, les rôles et les risques définis, l'étape suivante consiste à clarifier comment les décisions concernant l'IA sont prises en pratique. La gouvernance n'a pas besoin d'être complexe, mais elle doit être claire. Sans elle, les organisations font souvent face à une incertitude quant aux outils à utiliser, à la façon de réagir aux problèmes émergents ou à qui a l'autorité de prendre des décisions. À la base, cela signifie définir comment les nouveaux outils ou cas d'utilisation de l'IA sont évalués et approuvés. Par exemple, certaines organisations mettent en place un processus d'examen simple avant d'adopter un nouvel outil – tenant compte de facteurs tels que les pratiques de traitement des données, les risques potentiels et l'alignement avec la politique. Ce n'est pas nécessaire d'être formel ou long, mais cela devrait créer un moment de réflexion avant que les nouveaux outils ne soient largement utilisés.

Il est également important d'établir comment les préoccupations ou les problèmes sont soulevés et traités. Le personnel devrait savoir vers qui se tourner si quelque chose semble peu clair, risqué ou inapproprié, et se sentir à l'aise de poser des questions sans hésitation. Cela contribue à créer une culture où l'utilisation de l'IA est discutée ouvertement plutôt que traitée individuellement ou de manière informelle. Enfin, la gouvernance inclut la façon dont les décisions sont documentées et réévaluées. Même une documentation simple – comme noter pourquoi certains outils sont approuvés ou restreints – peut aider à assurer la cohérence dans le temps et à soutenir l'apprentissage à l'échelle de l'organisation. Comme pour le reste de la politique, la gouvernance devrait être perçue comme quelque chose qui évolue. Commencer par une structure simple et la raffiner au fur et à mesure que les usages se développent est souvent plus efficace que d'essayer de construire un système complet dès le départ.

Étape 8 : Soutenir le personnel par la formation et les orientations continues

Même la politique la plus solide ne suffira pas si le personnel ne sait pas comment l'appliquer en pratique. Les participantes et participants ont constamment souligné que les politiques d'IA doivent être soutenues par des orientations claires et continues – et pas seulement lors d'un déploiement ponctuel. Cela commence par l'établissement d'une compréhension de base partagée de ce que les outils d'IA peuvent et ne peuvent pas faire, y compris leurs limites, leurs risques et leurs cas d'utilisation appropriés. En pratique, cela signifie garder la formation simple, pertinente et ancrée dans le travail réel. Plutôt que de se concentrer sur la théorie, les organisations trouvent plus efficace d'utiliser des exemples concrets – comment l'IA pourrait être utilisée dans un rapport, une proposition ou une analyse de données – et de montrer ce qui est approprié, ce qui requiert de la prudence et ce qui devrait être évité. Créer un espace pour les questions et la discussion est tout aussi important, car le personnel rencontre souvent de l'incertitude en temps réel et a besoin d'un endroit pour l'explorer.

Le soutien continu est également important. À mesure que les outils évoluent et que les usages s'élargissent, le personnel continuera à rencontrer de nouvelles situations qui ne sont pas entièrement couvertes par la politique. Des points de contact réguliers – comme de courtes discussions lors des réunions d'équipe, de brefs rappels ou des exemples d'utilisation de l'IA partagés – contribuent à renforcer les attentes et à maintenir la politique active. Certaines organisations expérimentent également des approches légères, comme de courtes guides, des fiches de référence rapide ou des séances informelles, pour rendre le soutien plus accessible. En fin de compte, l'objectif est de renforcer la confiance et la cohérence. Lorsque le personnel comprend à la fois l'intention de la politique et la façon de l'appliquer dans son travail quotidien, l'utilisation de l'IA devient plus réfléchie, alignée et imputable.

Étape 9 : Réviser et adapter dans le temps

Les politiques d'IA sont les plus efficaces lorsqu'elles sont traitées comme des documents vivants plutôt que comme des règles fixes.

À mesure que les outils évoluent et que de nouveaux cas d'utilisation émergent, les politiques doivent s'adapter pour rester pertinentes et utiles. Les participantes et participants ont souligné que de nombreuses politiques deviennent rapidement obsolètes – non pas parce qu'elles sont mal conçues, mais parce qu'elles ne sont pas réévaluées. En pratique, cela signifie intégrer des points de révision réguliers. Certaines organisations fixent un cycle d'examen simple (par exemple, tous les 6 à 12 mois), tandis que d'autres révisent la politique de manière plus informelle en fonction des questions ou défis émergents. Ce qui importe le plus, c'est de créer une attente claire que la politique sera mise à jour dans le temps. Une façon pratique de soutenir cela est de s'appuyer sur des exemples concrets d'utilisation de l'IA au sein de l'organisation. Garder une trace de la façon dont l'IA est utilisée – ce qui fonctionne bien, là où surgit l'incertitude et là où des risques sont rencontrés – fournit une base concrète pour affiner la politique. Ces exemples contribuent à s'assurer que les mises à jour sont ancrées dans la pratique réelle plutôt que dans des suppositions. La révision continue est également l'occasion d'intégrer les commentaires du personnel et, le cas échéant, des partenaires ou des communautés. Cela renforce la propriété partagée et aide à s'assurer que la politique continue de refléter les réalités du travail. Plutôt que de viser une politique parfaite dès le départ, les organisations trouvent de la valeur à commencer par quelque chose de simple, à l'appliquer en pratique et à l'améliorer au fil du temps. Cette approche itérative permet aux politiques d'évoluer aux côtés de la technologie et des besoins de l'organisation.

Faire fonctionner votre politique d'IA en pratique : Conseils clés pour la mise en œuvre

Ces conseils sont tirés directement des discussions des participantes et participants et reflètent ce que les organisations trouvent efficace en pratique pour rendre les politiques d'IA utilisables et pertinentes.

- 1. L'intégrer dans le travail quotidien :** Intégrer les politiques dans les processus et outils existants. Plutôt que d'attendre du personnel qu'il consulte un document autonome, intégrez la politique dans les modèles de rapports et de propositions, les listes de vérification et les flux de travail, ou les processus d'analyse de données. Exemple : Ajouter une étape simple dans les modèles : ☐ IA utilisée ☐ Résultat vérifié pour l'exactitude ☐ Aucune donnée sensible incluse
- 2. Utiliser de vrais exemples pour rester pertinent·e :** Maintenir une liste partagée et évolutive de la façon dont l'IA est réellement utilisée dans l'équipe. Il pourrait s'agir d'un court document ou d'un fil Slack où le personnel note rapidement des exemples – à quoi ils ont utilisé l'IA, ce qui a bien fonctionné et ce qui semblait peu clair ou risqué (p. ex. « utile pour réfléchir à la structure d'un rapport », « nécessitait beaucoup de modifications », « pas à l'aise de téléverser des données d'entretien »). Revoir ces exemples concrets tous les quelques mois, donne à l'équipe une base pratique pour affiner la politique.
- 3. Être clair·e et précis·e :** Définir des cas d'utilisation spécifiques, pas seulement des principes généraux. Être explicite sur les limites des données dès le départ. Inclure une section « À quoi ça ressemble en pratique » directement dans la politique, avec des exemples clairs et concrets. Plutôt que de s'appuyer sur des formulations générales comme « protéger la confidentialité », les politiques devraient préciser exactement ce que cela signifie en pratique.
- 4. La garder courte et utilisable :** Rendre les politiques courtes, claires et faciles à utiliser. Les participantes et participants ont insisté sur l'importance de les rendre « aussi transparentes, claires et brèves que possible ». Se concentrer sur un langage concis, des orientations pratiques et supprimer le jargon inutile. Exemple : Créer une version d'une page à laquelle le personnel peut se référer rapidement dans son travail.
- 5. Créer une propriété partagée :** Certaines participantes et participants ont noté qu'impliquer toute l'équipe – même si c'est « un peu ennuyeux » – crée un réel sentiment d'appartenance et fait que la politique devient la responsabilité de toutes et tous.
- 6. Clarifier les rôles et l'imputabilité :** Des rôles clairs contribuent à s'assurer que les politiques restent actives et soutenues dans le temps. Définir les attentes envers le personnel, la direction et toute personne désignée (comme un·e champion·ne de la politique) rend les responsabilités visibles.

- 7. La garder visible et active :** Utiliser des points de contact simples et récurrents pour maintenir les politiques actives.
- Réviser une petite section de la politique (1 à 3 pages) lors des réunions d'équipe régulières. Exemple : Une invitation mensuelle de 5 minutes à l'équipe : « Quelqu'un a-t-il utilisé l'IA ce mois-ci d'une manière qui a soulevé des questions ? »
 - Faire tourner la responsabilité d'animer de courtes discussions entre les membres de l'équipe
 - Inclure de courts suivis (p. ex. quiz, réflexion, discussion)
 - Coupler les séances de révision de la politique avec des collations ou des activités interactives pour rendre l'engagement plus agréable
- 8. Relier la politique aux principes :** Les politiques sont souvent perçues comme restrictives ou non pertinentes. Présentez la politique comme aidant le personnel à prendre des décisions et à protéger les communautés et les données. La relier directement aux valeurs organisationnelles. Exemple : « Cette politique nous aide à nous assurer que notre travail reste exact, éthique et ancré dans les réalités des communautés. »
- 9. Commencer simple et construire dans le temps :** Les politiques n'ont pas besoin d'être parfaites dès le départ. Commencer par une approche simple et claire – et l'affiner au fil du temps en fonction de l'usage réel – réduit la pression et favorise l'apprentissage. Cela permet à la politique d'évoluer aux côtés de l'organisation et de la technologie.

Ressources clés en matière de politique d'IA

La séance a mis en lumière une gamme de ressources pratiques que les organisations peuvent utiliser pour développer ou affiner leurs politiques d'IA. Ce ne sont pas des solutions universelles, mais elles offrent de solides points de départ, des exemples et un langage pouvant être adaptés à votre contexte.

1. Principes et cadres fondamentaux

Ces ressources sont utiles lorsque vous définissez votre approche générale, vos valeurs et vos limites.

- **Principes de l'IA de l'OCDE :** Un cadre international largement utilisé énonçant les principes d'une IA digne de confiance et centrée sur l'humain. Utilisez-le pour ancrer votre politique dans des normes reconnues et des approches fondées sur les droits. <https://www.oecd.org/en/topics/sub-issues/ai-principles.html>
- **Fondation Mozilla :** Offre des perspectives pratiques centrées sur les personnes concernant l'IA, notamment des travaux sur l'IA responsable et la technologie d'intérêt public. Utilisez-la pour réfléchir de manière critique au pouvoir, à l'équité et aux impacts réels. <https://www.youtube.com/playlist?list=PLnRGhgZaGeBv4FyPOMMY3OOFQJAz8J0lb>

2. Exemples pratiques de politiques organisationnelles

Ces ressources sont utiles lorsque vous rédigez ou affinez votre propre politique.

- **Université de Toronto** – Lignes directrices sur l'utilisation de l'IA : Orientations claires sur l'utilisation acceptable, la confidentialité et les responsabilités du personnel. Utilisez-les pour des exemples concrets de structuration des orientations internes. <https://people.utoronto.ca/employees/administrative-staff-and-ai/guideline-on-the-use-of-artificial-intelligence-by-university-administrative-staff/>
- **University of British Columbia** – Principes pour l'utilisation de l'IA générative : Met l'accent sur l'expérimentation responsable et la mitigation des risques. Utilisez-les pour façonner des principes directeurs qui laissent encore de la flexibilité. <https://genai.ubc.ca/guidance/principles/> https://it-genai-2023.sites.olt.ubc.ca/files/2024/08/PRINCIPLES_FOR_MITIGATING_RISKS.pdf

- **Bibliothèque publique de Toronto – Politique sur l’IA** : Une politique pratique et opérationnelle axée sur l’utilisation des outils d’IA par le personnel. Utilisez-la comme modèle pour des politiques organisationnelles simples et applicables. <https://tpl.ca/policies-and-terms-of-use/artificial-intelligence-policy/>
- **Fondation Wikimedia – Stratégie IA** : Comprend un langage fort sur la façon dont l’utilisation de l’IA s’aligne sur la mission et les valeurs communautaires. Utilisez-la pour vous inspirer sur l’alignement de l’utilisation de l’IA avec l’objectif et l’imputabilité publique. https://meta.wikimedia.org/wiki/Strategy/Multigenerational/Artificial_intelligence_for_editors
- **Trousse d’outils de politique d’IA pour les OBNL canadiens – Mitchell Consulting Solutions** : Un guide complet, étape par étape, pour élaborer une politique d’utilisation de l’IA, spécialement adapté au secteur à but non lucratif. Couvre tout, des principes fondamentaux et des considérations éthiques à la gouvernance, la gestion des risques et la mise en œuvre. Utilisez-la si vous souhaitez une structure complète pour construire votre politique de zéro ou renforcer une politique existante. <https://mitchellconsultingsolutions.ca/resources/ai-policy-toolkit/#toolkit>

3. Orientations appliquées et sectorielles

Ces ressources montrent comment les organisations intègrent l’IA dans des contextes réels à enjeux élevés.

- **Groupe d’évaluation indépendante de la Banque mondiale – Note d’orientation sur l’IA** : Porte sur l’utilisation de l’IA dans l’évaluation tout en maintenant la rigueur et la qualité. Utilisez-la si vous travaillez avec des données, de la recherche ou de l’évaluation. https://ieg.worldbankgroup.org/sites/default/files/Data/Evaluation/files/report-artificial_intelligence_strategy.pdf
- **Stratégie IA de la Banque mondiale (IEG)** : Partage comment l’IA est introduite dans une organisation, y compris les limites et la portée. Utilisez-la pour réfléchir à une adoption progressive ou contrôlée. <https://documents1.worldbank.org/curated/en/099136005132515321/pdf/IDU-dccd6e52-4ee3-4294-a264-28fda8a94a49.pdf>
- **CICR – Politique sur l’intelligence artificielle (2024)** : Un solide exemple de politique d’IA centrée sur l’humain dans des environnements complexes et à risque élevé. Utilisez-la pour comprendre comment aborder l’IA là où les risques pour les personnes sont importants. <https://www.icrc.org/en/publication/building-responsible-humanitarian-approach-icrcs-policy-artificial-intelligence>

Réflexion de clôture

L’IA est déjà utilisée dans les organisations – souvent de manière informelle, inégale et sans orientation claire. Ce que cette séance a mis en évidence, c’est la nécessité de prendre des décisions intentionnelles concernant l’IA – si l’utiliser, comment l’utiliser et où ne pas l’utiliser – en se basant sur vos valeurs et votre contexte. Nous espérons que cette séance – et les réflexions et ressources qui y ont été partagées – vous sera utile alors que vous commencez ou poursuivez le développement de votre propre approche. Si vous réfléchissez ou travaillez à l’élaboration d’une politique d’IA, n’hésitez pas à nous contacter ou à poursuivre la conversation. C’est un travail en évolution, et il y a une réelle valeur à apprendre les un-e-s des autres au fil du chemin.

Il s’agit d’une série participative. Si vous avez des idées de thèmes à aborder, des personnes intervenantes potentielles, des études de cas pertinentes ou des questions à explorer, nous serions ravis de vous lire. Veuillez communiquer avec Paula Richardson à : richardson@salanga.org.