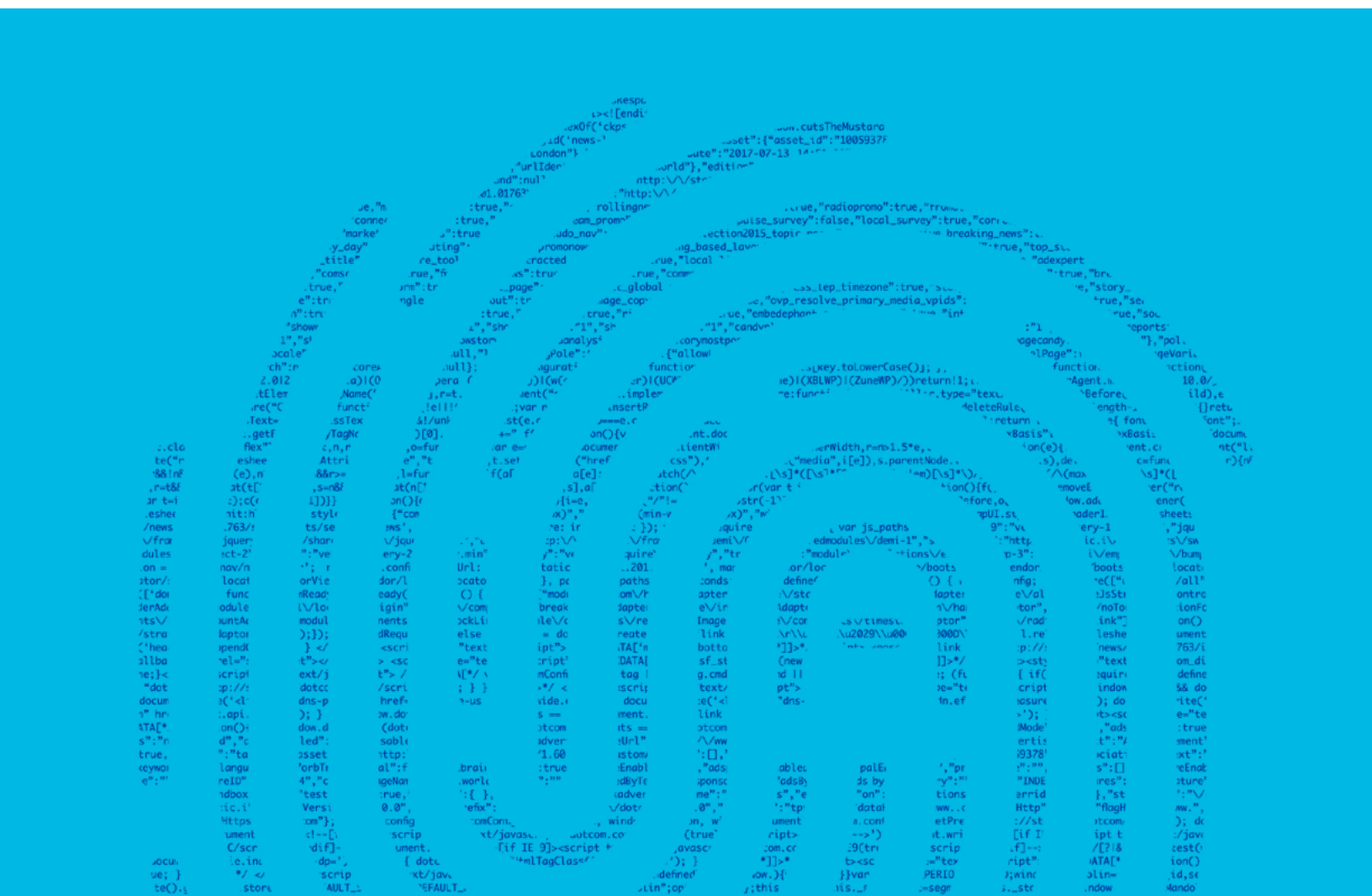




إعلان مونتريال

الذكاء الاصطناعي المسؤول_

إعلان مونتريال من أجل تنمية مسؤولة للذكاء الاصطناعي
لعام 2018



جدول المحتويات

تعد هذه الوثيقة جزءاً من إعلان مونتريال من أجل
تنمية مسؤولة للدَّكاء الاصطناعي لعام 2018.
يمكن العثور على التقرير الكامل على الرابط التالي :

5	قراءة الإعلان
7	التَّوطئة
المبادئ	
8	1. مبدأ الرِّفاه
9	2. مبدأ احترام الاستقلاليَّة
10	3. مبدأ حماية الخصوصية والحميميَّة
11	4. مبدأ التَّضامن
12	5. مبدأ المشاركة الديمقراطيَّة
13	6. مبدأ الإنصاف
14	7. مبدأ دمج التَّنوع
15	8. مبدأ الحيطة
16	9. مبدأ المسؤولية
17	10. مبدأ التَّمتية المستدامة
18	قاموس المصطلحات
I	الاعتمادات
II	الشُّركاء

تم اعتماد استخدام صيغة المذكر في هذه الوثيقة
لتسهيل القراءة، دون أي نية للتَّمييز.

قراءة إعلان مونتريال من أجل تنمية مسؤولة للذكاء الاصطناعي

ما الغرض من الإعلان؟

إعلان مونتريال من أجل تنمية مسؤولة للذكاء الاصطناعي ثلاثة أهداف رئيسية، وهي كما يلي:

1. وضع إطار عمل أخلاقي لتطوير الذكاء الاصطناعي ونشره؛

2. توجيه الانتقال الرقمي حتى يستفيد الجميع من هذه الثورة التقنية؛

3. فتح فضاء نقاش وطني ودولي لتحقيق تنمية شاملة ومنصفة ومستدامة بيئيًا للذكاء الاصطناعي.

عمّ يدور الإعلان؟

المبادئ

يتمثل الهدف الأول من الإعلان في تحديد المبادئ والقيم الأخلاقية التي تعزز المصالح الأساسية للأفراد والمجموعات. تبقى هذه المبادئ السارية على المجال الرقمي و الذكاء الاصطناعي عامة ومجرّدة. من المهم أخذ النقاط التالية بعين الاعتبار لقراءتها بشكل صحيح:

< على الرغم من تقديمها في شكل قائمة، فإنه لا يوجد تسلسل هرمي، إذ لا يقل المبدأ الأخير أهمية عن المبدأ الأول. لكن من الممكن إيلاء أهمية لأحد المبادئ أكثر من غيره أو اعتبار أحدها أكثر ملائمة من غيره، وفقًا للظروف.

< على الرغم من تنوعها، فإنه يجب أن يتم تأويلها بشكل متناسق لمنع حدوث أي تعارض من شأنه أن يمنع تطبيقها. ترسم حدود تطبيق أحد المبادئ، كقاعدة عامة، من خلال مجال تطبيق مبدأ آخر.

< على الرغم من أنها تعكس الثقافة الأخلاقية والسياسية للمجتمع الذي احتضن إعدادها فإنها تمثل أساسا للحوار الثقافي والدولي.

< على الرغم من إمكانية تأويلها بطرق مختلفة، فإنه لا يمكن تأويلها كيفما اتفق. من الضروري أن يكون التّأويل متناسقا.

< على الرغم من كونها مبادئ أخلاقية، فإنه يمكن ترجمتها إلى لغة سياسية وتفسيرها بأسلوب قانوني.

من هذه المبادئ تم إعداد توصيات تهدف إلى اقتراح إرشادات توجيهية لتحقيق الانتقال الرقمي ضمن الإطار الأخلاقي للإعلان. وهي تهدف إلى تغطية بعض الموضوعات الرئيسية الشاملة العابرة للقطاعات وذلك للتفكير في الانتقال نحو مجتمع يُمكن فيه الذكاء الاصطناعي من تعزيز المصلحة العامة: الإدارة الخوارزمية ومحو الأمية الرقمية والدمج الرقمي للتنوع والاستدامة البيئية.

لمن يتوجه الإعلان؟

يتوجه إعلان مونتريال لأي شخص أو منظمة من المجتمع المدني أو شركة ترغب في المشاركة في التطوير المسؤول للذكاء الاصطناعي، سواء كان ذلك من أجل المساهمة علمياً أو تقنياً، أو لتطوير مشاريع اجتماعية، أو لوضع قواعد (لوائح وقوانين) تنطبق عليها أو للتمكّن من معارضة مناهج سيئة أو غير حكيمة، أو كي تستطيع تنبيه الرأي العام عند الضرورة.

كما أنه يتوجه إلى الممثلين السياسيين المنتخبين والمعيينين، والذين يتوقع منهم مواطنوهم

تقييم التغييرات الاجتماعية الكامنة ووضع إطار عمل يمكن من الانتقال الرقمي، على وجه السرعة، خدمة للمصلحة العامة.

كما يتوقع هؤلاء المواطنون منهم استباق المخاطر التي يُمثلها تطوير الذكاء الاصطناعي.

إعلان وفقاً لأي منهج ؟

نابع الإعلان من عملية تداول شاملة ضمت مواطنين وخبراء وسلطات عمومية وأطراف صناعية فاعلة ومنظمات المجتمع المدني وهيئات مهنية. ولهذا النهج أهمية ثلاثية:

1. التحكيم الجماعي في الخلافات الأخلاقية والمجتمعية حول الذكاء الاصطناعي؛

2. تحسين جودة التفكير في الذكاء الاصطناعي المسؤول؛

3. تعزيز شرعية المقترحات من أجل ذكاء اصطناعي مسؤول.

يعد وضع المبادئ والتوصيات عملاً مشتركاً شمل مجموعة متنوعة من المشاركين في الفضاءات العامة وقاعات اجتماعات المؤسسات المهنية وحول الموائد المستديرة لخبراء دوليين وفي مكاتب بحوث وفصول دراسية وعبر الإنترنت، مع احترام نفس المستوى من الصرامة باستمرار.

ماذا بعد الإعلان؟

نظراً لأن الإعلان يهتم بتقنية ما انفكت تتطور بثبات منذ الخمسينيات، تقنية ما فتئت وتيرة ابتكاراتها الرئيسية تتسارع بنسق أسي، فإنه من الضروري اعتبار الإعلان وثيقة توجيهية مفتوحة يمكن مراجعتها ومواءمتها وفقاً لتطور المعرفة والتقنيات، إلى جانب ردود الفعل على استخدام الذكاء الاصطناعي في المجتمع. وقد وصلنا في نهاية عملية إعداد الإعلان إلى منطلق محادثة مفتوحة ومدمجة حول مستقبل البشرية الذي تخدمه تقنيات الذكاء الاصطناعي.

لأول مرة في تاريخ البشرية يمكن إنشاء أنظمة مستقلة قادرة على إنجاز المهام المعقدة التي كان يُعتقد أنها حصرية للذكاء الطبيعي، مثل: معالجة كميات كبيرة من المعلومات والحساب والتنبؤ والتعلم وتكييف الإجابات مع الوضعيات المتغيرة، إلى جانب التعرف على الأشياء وتصنيفها.

نظراً للطبيعة غير المادية لهذه المهام، وقياساً على الذكاء البشري، فإننا نطلق على هذه الأنظمة المتنوعة مصطلحاً عاماً وهو الذكاء الاصطناعي. يُشكل الذكاء الاصطناعي نموذجاً رئيسياً للتقدم العلمي والتكنولوجي، والذي يمكن أن تنتج عنه فوائد اجتماعية كبيرة من خلال تحسين ظروف المعيشة والصحة وتيسير العدالة وخلق الثروات وتعزيز السلامة العامة والتحكم في تأثير الأنشطة البشرية على البيئة والمناخ.

ولا تكفي الآلات الذكية بإنجاز العمليات الحسابية أفضل من البشر بل يمكنها أيضاً التفاعل مع الكائنات الحية ومرافقتها والعناية بها.

ومع ذلك، يواجه تطور الذكاء الاصطناعي تحديات أخلاقية ومخاطر اجتماعية كبرى.

يمكن للآلات الذكية أن تقوم في الواقع بتقييد اختيارات الأفراد والمجموعات والتخفيض من جودة الحياة وقلب تنظيم العمل وسوق الشغل والتأثير على الحياة السياسية والتعارض مع الحقوق الأساسية. كما يمكن لها أن تساهم في تفاقم عدم المساواة الاقتصادية والاجتماعية والتأثير في الأنظمة الإيكولوجية والبيئة والمناخ.

وعلى الرغم من عدم وجود تقدم علمي أو حياة اجتماعية خالية من المخاطر، فإن مهمة تحديد الغايات الأخلاقية والسياسية التي تعطي معنى لما نواجهه من مخاطر في عالم يخلو من اليقين تعود إلى المواطنين.

ستكون فوائد الذكاء الاصطناعي أكبر كلما قلت المخاطر المرتبطة بنشره. ويتمثل الخطر الأول لتطوير الذكاء الاصطناعي في الإيهام بأنه يمكننا السيطرة على المستقبل من خلال الحسابات.

وبشكل اختزال المجتمع في سلسلة من الأرقام وحكمه من خلال الإجراءات الخوارزمية حلماً قديماً لا يزال يغذي طموحات الإنسان.

إلا أن الغد نادراً ما يشبه اليوم، كما أن الأرقام تعجز عن تحديد ما الذي له قيمة أخلاقية أو

ما هو مرغوب فيه اجتماعياً، عندما يتعلق الأمر بالشؤون الإنسانية.

تُعد مبادئ هذا الإعلان بوصلة أخلاقية تمكن من توجيه تطوير الذكاء الاصطناعي نحو الغايات الأخلاقية والاجتماعية المنشودة.

كما أنها توفر إطاراً أخلاقياً يعزز حقوق الإنسان المعترف بها دولياً في المجالات المعنية بتطبيق الذكاء الاصطناعي.

بشكل عام، تضع المبادئ المذكورة الأساس لتنمية الثقة الاجتماعية تجاه أنظمة الذكاء الاصطناعي.

هذا وتستند مبادئ هذا الإعلان إلى الاعتقاد المشترك بأن البشر يسعون إلى الانتعاش ككائنات اجتماعية تتمتع بالأحاسيس والمشاعر والأفكار، وأنهم يسعون جاهدين لتحقيق إمكاناتهم الكامنة من خلال ممارسة قدراتهم العاطفية والأخلاقية والفكرية بحرية.

ويقع على عاتق مختلف الجهات الفاعلة والصناعة للقرار، العامة والخاصة،

سواء كان ذلك على المستوى المحلي أو الوطني أو الدولي، ضمان توافق

تطوير ونشر الذكاء الاصطناعي مع حماية وانتعاش القدرات البشرية الأساسية.

وسعيًا إلى بلوغ هذا الهدف، يجب تأويل المبادئ المقترحة بطريقة متناسقة، مع

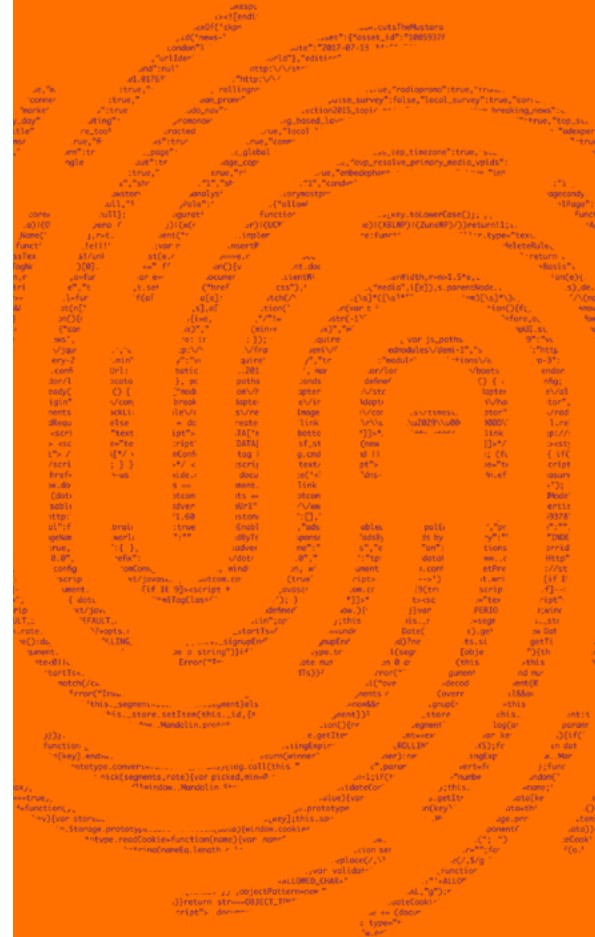
مراعاة خصوصية السباقات الاجتماعية والثقافية والسياسية والقانونية لتطبيقها.

مبدأ الرفاه

1

يجب أن يسمح تطوير أنظمة الذكاء الاصطناعي (SIA) واستخدامها بتنمية رفاهية جميع الكائنات الحية.

1. يجب أن تساعد أنظمة الذكاء الاصطناعي الأفراد على تحسين ظروفهم المعيشية وصحتهم وظروف عملهم.
2. يجب أن تسمح أنظمة الذكاء الاصطناعي للأفراد بإرضاء خياراتهم طالما أنهم لا يسببون أي ضرر لغيرهم من الكائنات الحية.
3. يجب أن تسمح أنظمة الذكاء الاصطناعي للأفراد بممارسة قدراتهم الجسدية والعقلية.
4. يجب ألا تمثل أنظمة الذكاء الاصطناعي مصدرا للشعور بالضيق إلا إذا مكن هذا الشعور من توليد رفاهية كبرى لا يمكن الوصول إليها بطريقة أخرى.
5. يجب ألا يساهم استخدام أنظمة الذكاء الاصطناعي في زيادة التوتر أو القلق أو الشعور بنوع من الحرش له علاقة بالبيئة الرقمية.



مبدأ احترام الإستقلالية

2

1. يجب أن تمكّن أنظمة الذكاء الاصطناعي الأفراد من تحقيق أهدافهم الأخلاقية ومفهومهم للحياة التي تستحق العيش.
2. يجب ألا تطوّر أنظمة الذكاء الاصطناعي أو تستخدم لفرض نمط حياة معين على الأفراد، سواء كان ذلك بطريقة مباشرة أو غير مباشرة، من خلال وضع التّيات إكراهية للمراقبة والتّقييم والتّحفيز.
3. يجب ألا تستخدم المؤسسات العامة أنظمة الذكاء الاصطناعي لنشر أو تشويه مفهوم معيّن للحياة الجيدة.
4. من الضروري تنمية إمكانات المواطنين فيما يتعلق بالتّقنيات الرّقمية وذلك من خلال ضمان الوصول إلى مختلف أنواع المعرفة المناسبة وتنمية المهارات الهيكلية (محو الأمية الرّقمية والإعلامية) وتشكيل التفكير النّقدي.
5. يجب ألا تطوّر أنظمة الذكاء الاصطناعي لنشر معلومات غير موثوقة أو أكاذيب أو دعايات، بل يجب على العكس تصميمها بهدف تقليص انتشارها.
6. يجب أن يجنّب تطوير أنظمة الذكاء الاصطناعي خلق التّبعية عبر تقنيات جاذبة للانتباه وعبر محاكاة الصفات البشرية بطرق من شأنها أن تسبب التّبايناً بين أنظمة الذكاء الاصطناعي والبشر.

يجب تطوير أنظمة الذكاء الاصطناعي واستخدامها في ظلّ مراعاة استقلالية الأفراد، وبغاية تنمية تحكّم الأفراد في حياتهم وبينتهم.



مبدأ حماية الخصوصية والسرية

يجب حماية الخصوصية
والسرية من تدخل أنظمة
الدكاء الاصطناعي وأنظمة
التقاط وأرشفة البيانات
الشخصية.

1. يجب حماية المساحات الحميمة التي لا يخضع فيها الأشخاص للمراقبة أو التقييم الرقمي من تدخل أنظمة الدكاء الاصطناعي وأنظمة التقاط وأرشفة البيانات الشخصية.
2. يجب حماية حميمية الأفكار والعواطف من استخدامات أنظمة الدكاء الاصطناعي وأنظمة التقاط وأرشفة البيانات الشخصية القادرة على إلحاق الضرر، لاسيما الاستخدامات التي تهدف إلى إطلاق أحكام أخلاقية على خيارات الأشخاص أو على خياراتهم الحياتية.
3. يجب أن تتاح للأشخاص بصفة دائمة إمكانية اختيار قطع الاتصال الرقمي في حياتهم الخاصة، ويجب أن توفر أنظمة الدكاء الاصطناعي بشكل صريح خيار قطع الاتصال في فترات زمنية منتظمة، دون تشجيع الأشخاص على البقاء متصلين.
4. يجب أن يتمتع الأفراد بتحكم واسع في المعلومات المتعلقة بخياراتهم. يجب ألا تنشئ أنظمة الدكاء الاصطناعي ملامح اختيارات فردية للتأثير على سلوك الأفراد المعنيين دون موافقتهم الحرة والمستنيرة.
5. يجب أن تضمن أنظمة التقاط وأرشفة البيانات الشخصية سرية البيانات وإخفاء هوية ملفات التعريف الشخصية.
6. يجب أن يتاح لكل شخص التحكم الواسع في بياناته الشخصية، ولاسيما في ما يتعلق بجمعها واستخدامها ونشرها. ولا يمكن أن يكون استعمال الأشخاص لأنظمة الدكاء الاصطناعي والخدمات الرقمية مشروطاً بتخليهم عن ملكية بياناتهم الشخصية.
7. يمكن لكل شخص أن يتبرع ببياناته الشخصية للمؤسسات البحثية للمساهمة في التقدم بالمعرفة.
8. يجب ضمان سلامة الهوية الشخصية. يجب ألا تستخدم أنظمة الدكاء الاصطناعي لمحاكاة مظهر أحد الأشخاص، أو صوته أو أي صفات فردية أخرى أو تغييرها بهدف الإضرار بسمعته أو التلاعب بأشخاص آخرين.

مبدأ التضامن

4

يجب أن يتوافق تطوير
أنظمة الذكاء الاصطناعي مع
الحفاظ على أواصر التضامن
بين الأشخاص والأجيال.

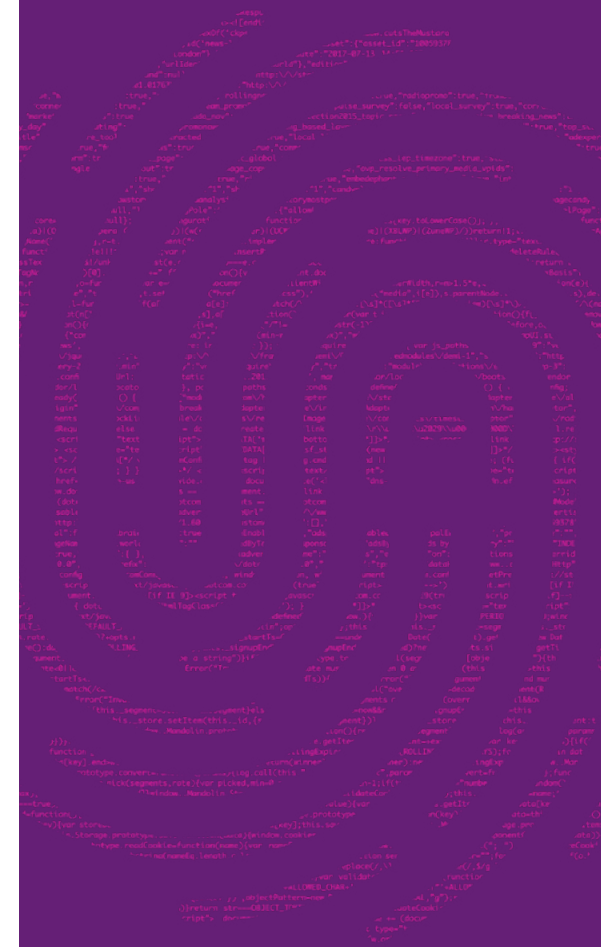
1. يجب ألا تهدد أنظمة الذكاء الاصطناعي الحفاظ على العلاقات الإنسانية العاطفية والمعنوية، ويجب أن يقع تطويرها بغاية تعزيز هذه العلاقات والحد من هشاشة الأشخاص وعزلتهم.
2. يجب أن تطور أنظمة الذكاء الاصطناعي بهدف التعاون مع البشر في المهام المعقدة، ويجب أن تعزز العمل التعاوني بين البشر.
3. ينبغي ألا يقع وضع أنظمة الذكاء الاصطناعي لاستبدال الأشخاص في مهام تتطلب علاقات إنسانية ذات جودة، بل ينبغي تطويرها لتيسير هذه العلاقات.
4. ينبغي أن تأخذ أنظمة الرعاية الصحية التي تستخدم أنظمة الذكاء الاصطناعي بعين الاعتبار أهمية العلاقات مع الأسرة وموظفي الرعاية الصحية بالنسبة للمرضى.
5. يجب ألا يحدث تطوير أنظمة الذكاء الاصطناعي على تبني سلوك قاس تجاه الروبوتات المصممة لمحاكاة مظهر البشر أو الحيوانات والتي تتصرف مثلهم ظاهرياً.
6. يجب أن تمكن أنظمة الذكاء الاصطناعي من تحسين إدارة المخاطر وتهيئة الظروف لبناء مجتمع أكثر فعالية يقوم على التشارك في تحمل المخاطر الفردية والجماعية.



مبدأ المشاركة الديمقراطية

ينبغي أن تفي أنظمة الذكاء الاصطناعي بمعايير الوضوح والقدرة على التبرير والوصول، ويجب أن تكون قابلة للخضوع للمراجعة والنقاش والمراقبة الديمقراطية.

1. يجب أن يكون سير عمل أنظمة الذكاء الاصطناعي المعنية باتخاذ قرارات مؤثرة على حياة الأشخاص أو جودة حياتهم أو سمعتهم، مفهوما لمبتكره.
2. يجب أن تكون القرارات التي تتخذها أنظمة الذكاء الاصطناعي والتي تؤثر على حياة الأشخاص أو جودة حياتهم أو سمعتهم معللة دائماً بلغة يفهمها مستخدموها أو الأشخاص الذين يعيشون عواقب استخدامها.
- ويتمثل التعليل في تقديم أهم عوامل ومعايير أخذ القرار، ويجب أن يكون مشابهاً للتعليلات التي قد نطالب بها إنساناً يتخذ نفس نوعية القرار.
3. يجب أن يكون استخدام شفرة الخوارزميات العامة أو الخاصة في متناول السلطات العمومية المختصة وأصحاب المصلحة المعنيين بالأمر وذلك بهدف التحقق والمراقبة.
4. يجب الإبلاغ عن اكتشاف أي أخطاء متعلقة بتشغيل أنظمة الذكاء الاصطناعي أو بأعراض غير متوقعة أو غير مرغوب فيها أو بخروقات أمنية أو بتسريبات بيانات، إلى الجهات العمومية المختصة وأصحاب المصلحة المعنيين بالأمر والأشخاص المتضررين من الوضع.
5. وفقاً لمتطلبات الشفافية في القرارات العامة، يجب أن تكون شفرة خوارزميات اتخاذ القرارات التي تستخدمها السلطات العمومية في متناول الجميع، باستثناء الخوارزميات التي تمثل خطراً جدياً ذي نسبة احتمال مرتفعة في حالة إساءة استخدامها.
6. في ما يخص أنظمة الذكاء الاصطناعي العامة ذات التأثير الهام على حياة المواطنين، فإنه يجب أن تتوفر لهؤلاء الفرصة والكفاءة لمناقشة المعايير الاجتماعية لأنظمة الذكاء الاصطناعي وأهدافها وحدود استعمالها.
7. يجب أن نكون قادرين على التحقق في أي وقت من أن أنظمة الذكاء الاصطناعي تقوم بما يُرجى من أجله وتستخدم في نفس الإطار.
8. يجب أن يكون أي مستخدم على علم بأي قرار متخذ من قبل نظام الذكاء الاصطناعي، قد يتعلّق به أو يمسّه.
9. يجب أن يتمكن أي مستخدم لخدمة تلجى إلى استعمال عميل محادثات، من تحديد ما إذا كان يتفاعل مع نظام الذكاء الاصطناعي أو مع شخص حقيقي، بسهولة.
10. يجب أن تظل البحوث في مجال الذكاء الاصطناعي مفتوحة وفي متناول الجميع.

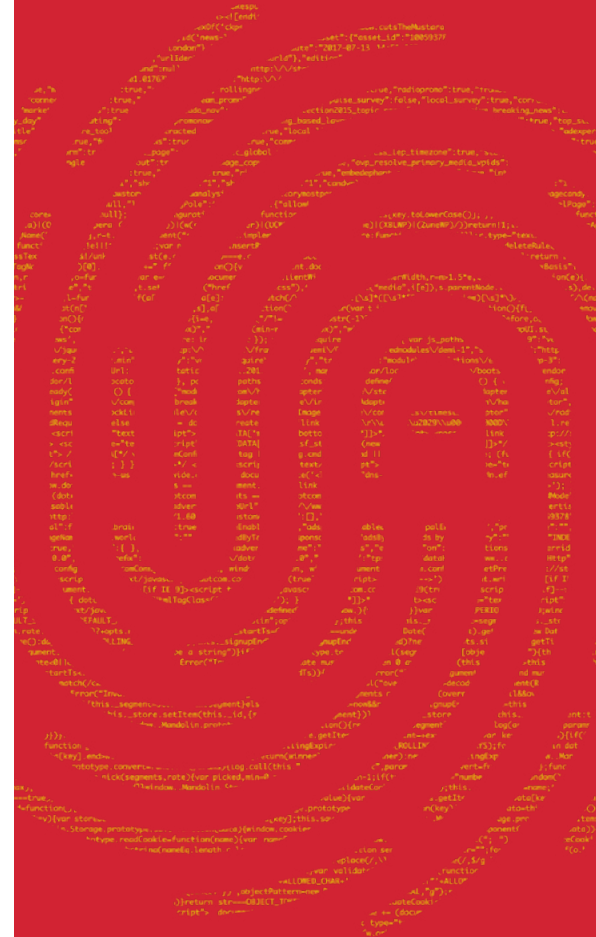


مبدأ الإنصاف

6

1. وجب وضع أنظمة الذكاء الاصطناعي وتدريبها على ألا تخلق أو تعزز أو تعيد إنتاج التمييز القائم على أساس الاختلافات الاجتماعية و الجنسية و العرقية و الثقافية و الدينية وغيرها.
2. يجب أن يساهم تطوير أنظمة الذكاء الاصطناعي في القضاء على علاقات الهيمنة القائمة على التفاوت في السلطة أو الثروة أو المعرفة بين الأشخاص والمجموعات.
3. يجب أن تتمحور عن تطوير أنظمة الذكاء الاصطناعي فوائد اقتصادية واجتماعية للجميع وذلك عبر الحد من أوجه عدم المساواة والهشاشة الاجتماعيتين.
4. يجب أن يتوافق التطوير الصناعي لأنظمة الذكاء الاصطناعي مع ظروف العمل اللائقة، في كل مرحلة من دورة حياة هذه الأنظمة، بدءًا من استخراج الموارد الطبيعية إلى إعادة تدويرها، ومرورًا بمعالجة البيانات.
5. يجب اعتبار النشاط الرقمي لمستخدمي أنظمة الذكاء الاصطناعي والخدمات الرقمية شغلًا يساهم في عمل الخوارزميات و إنتاج القيمة.
6. يجب ضمان قدرة الجميع على الوصول إلى الموارد والمعرفة والأدوات الرقمية الأساسية.
7. يمثل تطوير الخوارزميات العامة والبيانات المفتوحة، من أجل تدريبها وتشغيلها، هدفًا منصفًا اجتماعيًا وجب دعمه.

يجب أن يساهم تطوير
أنظمة الذكاء الاصطناعي
واستخدامها في إنشاء
مجتمع عادل ومنصف.



مبدأ دمج التنوع

7

1. يجب ألا يؤدي تطوير نظام الذكاء الاصطناعي واستخدامه إلى تجانس المجتمع من خلال توحيد السلوكيات والآراء.

2. يجب أن يأخذ تطوير نظام الذكاء الاصطناعي بعين الاعتبار مختلف تعبيرات التنوع الاجتماعي والثقافي من اللحظة التي يتم فيها تصميم الخوارزميات.

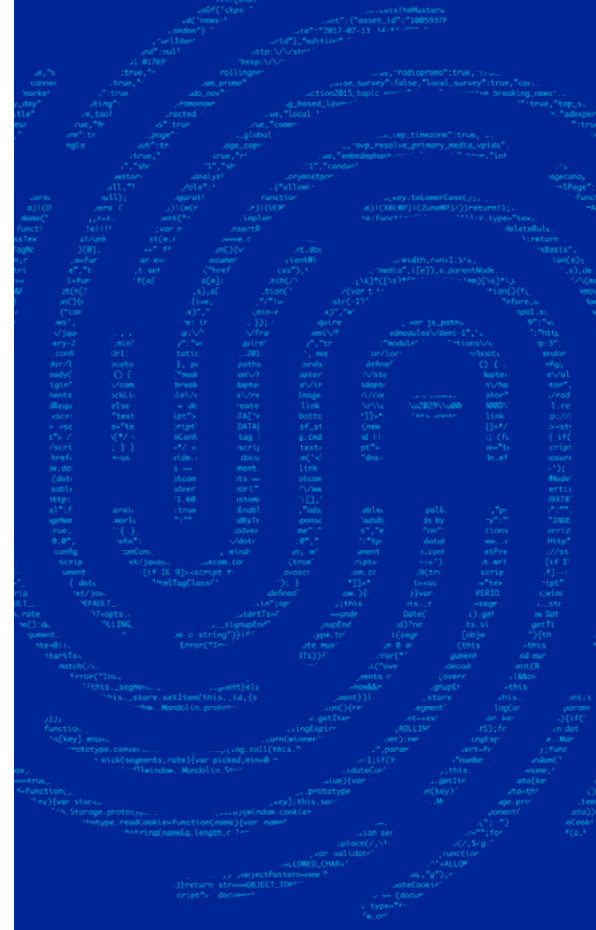
3. يجب أن تتسم أوساط تطوير الذكاء الاصطناعي بالشمولية وأن تعكس تنوع الأفراد والجماعات في المجتمع، سواء كان ذلك في مجال البحث أو الصناعة.

4. يجب على أنظمة الذكاء الاصطناعي تجنب حصر الأشخاص في ملف تعريف المستخدم أو في فقاعة مصفية، كما يجب ألا تقوم بتثبيت الهوية الشخصية عبر معالجة بيانات أنشطة الأشخاص السابقة وألا تقوم بتقليص خيارات التطوير الشخصية الخاصة بهم، لاسيما في مجال التعليم والعدالة والأعمال التجارية.

5. يجب ألا تطوّر أنظمة الذكاء الاصطناعي أو تستخدم بهدف الحد من حرية التعبير عن الأفكار أو الآراء والتي يعد تنوعها شرطاً لإقامة مجتمع ديمقراطي.

6. يجب تنويع عروض أنظمة الذكاء الاصطناعي، لكل فئة من الخدمات، لمنع الاحتكار الفعلي من التشكل وإلحاق الضرر بالحريات الفردية.

يجب أن يتوافق تطوير نظام الذكاء الاصطناعي واستخدامه مع الحفاظ على التنوع الاجتماعي والثقافي، كما يجب ألا يقيّد نطاق الخيارات الحياتية أو التجارب الشخصية.



مبدأ الحيطة

8

1. من الضروري تطوير آليات تراعي إمكانيات الاستخدام المزدوج، المفيدة والصّارة، لأبحاث الذكاء الاصطناعيّ (سواء كانت عامّة أو خاصة) ولتطوير أنظمة الذكاء الاصطناعيّ، وذلك للحد من الاستخدامات الصّارة.

2. عندما تُعرّض إساءة استخدام نظام الذكاء الاصطناعيّ، بنسبة احتمال مرتفعة، السلامة أو الصّحة العامّة للخطر، فإنه من الحكمة تقليص انتشار خوارزمياته والوصول المفتوح إليها.

3. يجب أن تفي أنظمة الذكاء الاصطناعيّ بمعايير الموثوقيّة والأمان والنّزاهة، وأن تخضع لاختبارات لا تُعرّض حياة الأشخاص للخطر أو تُضرّ بجودة حياتهم أو تؤثر سلبيّا على سمعتهم أو سلامتهم النفسيّة وذلك قبل طرحها في السّوق، سواء كانت معروضة مقابل رسوم أو مجانًا. ويجب أن تكون هذه الاختبارات مفتوحة أمام السلطات العموميّة المختصة وأصحاب المصلحة المعنيّين بالأمر.

4. يجب أن يقي تطوير أنظمة الذكاء الاصطناعيّ من مخاطر إساءة استخدام بيانات المستخدم وحماية سلامة البيانات الشّخصيّة وسريّتها.

5. يجب أن تنشر الأخطاء والخروقات المكتشفة في أنظمة الذكاء الاصطناعيّ وأنظمة التّقاط وأرشفة البيانات الشّخصيّة علنًا، وعلى نطاق عالمي، من قبل المؤسسات العموميّة والشركات، وذلك في القطاعات التي تُشكّل خطرًا كبيرًا على السلامة الشّخصيّة والتنظيم الاجتماعيّ.

يجب أن يتوخى كلّ من يشارك في تطوير أنظمة الذكاء الاصطناعيّ الحيطة من خلال استباق النّتائج السّلبية لاستخدام أنظمة الذكاء الاصطناعيّ بقدر الإمكان، واتّخاذ التدابير المناسبة لتلافيها.



مبدأ المسؤولية

1. يتحمل البشر وحدهم مسؤولية القرارات النابعة من التوصيات التي تقدمها أنظمة الذكاء الاصطناعي والإجراءات الناتجة عنها.
2. يجب أن يعود القرار النهائي إلى الإنسان في جميع المجالات التي يتعين فيها أخذ قرار من شأنه التأثير في حياة شخص ما أو في سمعته أو في جودة حياته، وينبغي أن يكون هذا القرار حراً ومستقراً.
3. يجب أن يعود أي قرار بالقتل دائماً للبشر ولا يمكن أن تنقل مسؤولية اتخاذ هذا القرار إلى أحد أنظمة الذكاء الاصطناعي.
4. يتحمل الأشخاص الذين يعطون الإذن لأنظمة الذكاء الاصطناعي بارتكاب جريمة أو جنحة أو الذين يسمحون لها بارتكابها بسبب إهمالهم، مسؤولية هذه الجريمة أو الجنحة.
5. عندما يتسبب نظام الذكاء الاصطناعي في حصول ضرر أو تلف ما، ويثبت أنه موثق وأن استخدامه تم بطريقة طبيعية فإنه من غير المعقول إلقاء اللوم على الأشخاص المشاركين في تطويره أو استخدامه.

يجب ألا يساهم تطوير أنظمة الذكاء الاصطناعي واستخدامها في إلغاء مسؤولية البشر عندما ينبغي اتخاذ قرار ما.



مبدأ التنمية المستدامة

10

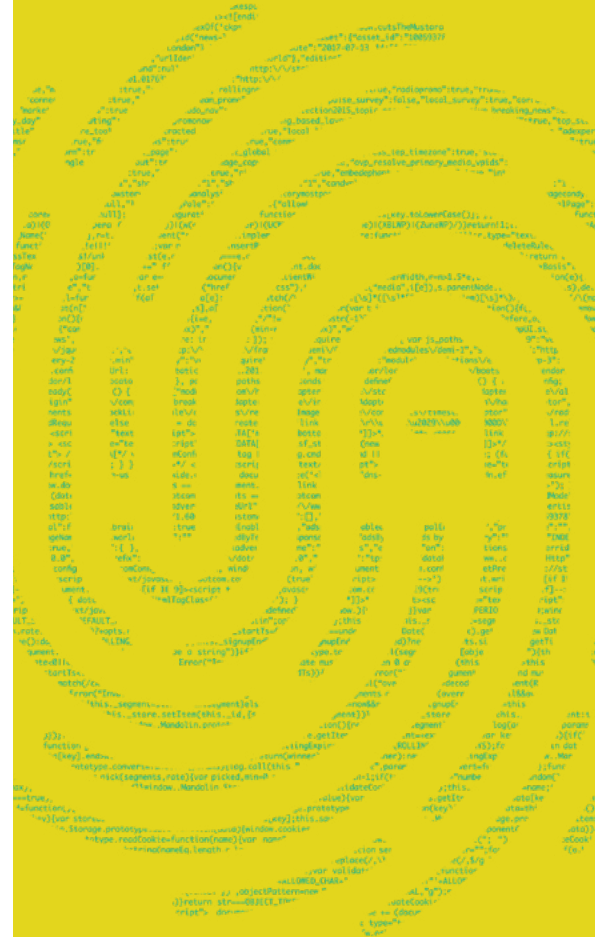
يجب تطوير أنظمة الذكاء الاصطناعي واستخدامها بطريقة تضمن استدامة بيئية قوية لكوكب الأرض.

1. يجب أن تهدف أجهزة أنظمة الذكاء الاصطناعي وبنيتها التحتية الرقمية وكل الأشياء المتصلة التي تعتمد عليها، مثل مراكز البيانات، إلى تحقيق أكثر فعالية ممكنة في استخدام الطاقة والحد من انبعاثات الغازات الدفينة وذلك على مدار دورة حياتها.

2. يجب أن تهدف أجهزة أنظمة الذكاء الاصطناعي وبنيتها التحتية الرقمية وكل الأشياء المتصلة التي تعتمد عليها، إلى توليد أقل قدر ممكن من النفايات الكهربائية والإلكترونية، وتوفير قنوات للصيانة والإصلاح وإعادة التدوير وفقًا لمبادئ الاقتصاد الدائري.

3. يجب أن تقلل أجهزة أنظمة الذكاء الاصطناعي وبنيتها التحتية الرقمية وكل الأشياء المتصلة التي تعتمد عليها من تأثيرنا في النظم الإيكولوجية والتنوع البيولوجي، في كل مرحلة من مراحل دورة حياتها، بما في ذلك خلال استخراج الموارد الطبيعية ومراحل نفاذ فوائدها.

4. يجب أن تدعم كل من الجهات الفاعلة العامة والخاصة تطوير أنظمة الذكاء الاصطناعي المسؤولة بيئيًا من أجل مكافحة إهدار الموارد الطبيعية والسلع المنتجة، ووضع سلاسل إمداد ومبادلات تجارية مستدامة، إلى جانب الحد من التلوث على المستوى العالمي.



الاستدامة البيئية القوية

يشير مفهوم الاستدامة البيئية القوية إلى فكرة وجوب توافق وتيرة استهلاك الموارد الطبيعية والانبعاث الملوثة مع الحدود البيئية لكوكبنا وتيرة تجدد الموارد والأنظمة البيئية، إلى جانب استقرار المناخ، وذلك لكي تكون مستدامة. وعلى عكس الاستدامة الضعيفة والأقل تطلّبا، لا تسمح الاستدامة القوية باستبدال خسارة الموارد الطبيعية من أجل رأس مال صناعي.

البيانات الشخصية

البيانات الشخصية هي تلك التي تمكّن من التعرف على شخص ما، سواء كان ذلك بطريقة مباشرة أو غير مباشرة.

البيانات المفتوحة

تمثّل البيانات المفتوحة بيانات رقمية يمكن للمستخدمين الوصول إليها بحريّة، مثل معظم نتائج الأبحاث المنشورة في الذكاء الاصطناعيّ.

التأثير الارتدادي

يمثل التأثير الارتدادي الآلية التي تؤدي من خلالها زيادة النجاعة الطاقية أو تحسّن الأداء البيئي للمنتجات والمعدات والخدمات إلى زيادة نسبية مرتفعة في استهلاكها. على سبيل المثال، تنتج بشكل عام عن ازدياد حجم الشاشات وارتفاع عدد الأجهزة الإلكترونية في البيوت وقطع مسافات أطول، سواء كان ذلك بالسيارة أو بالطائرة، زيادة في الضّغط على الموارد والبيئة.

التبعية للمسار

تمثّل التبعية للمسار آلية اجتماعية يمكن من خلالها اتخاذ قرارات تكنولوجية أو تنظيمية أو مؤسّساتية كانت تعتبر عقلانية يوما ما ولكنها أصبحت دون المستوى الأمثل اليوم، والتي تستمرّ رغم ذلك في التأثير على صنع القرارات. ويتم الحفاظ على هذه الآلية بسبب التّحيز المعرفي أو لأن تغييرها سيتطلّب تكلفة أو جهدا كبيرين للغاية وينطبق ذلك على البنية التحتية للطرق الحضرية عندما تؤدّي إلى إنشاء برامج تحسين حركة المرور عوضا عن طرح تغيير ينظّم التّنقل بشكل يخفّض من انبعاثات الكربون. ويجب أن تكون هذه الآلية معروفة عند استخدام الذكاء الاصطناعيّ في المشاريع الاجتماعية إذ قد تعزز بيانات التدريب في التعلّم الخاضع للإشراف أحيانا النماذج التنظيمية القديمة التي تكون جدواها قد أصبحت محل نقاش في الوقت الراهن.

التدريب

التدريب هو عملية التعلّم الآليّ التي يُصمّم نظام الذكاء الاصطناعيّ من خلالها نموذجًا من البيانات. ويتوقف أداء نظام الذكاء الاصطناعيّ على جودة النموذج الذي يتوقّف بدوره على كميّة وجودة البيانات المستخدمة أثناء التدريب.

التعلّم الآليّ

التعلّم الآليّ هو أحد فروع الذكاء الاصطناعيّ الذي يهتم بتصميم خوارزميات تسمح للنظام بالتعلّم من تلقاء نفسه. من بين مختلف التقنيات، نستطيع تمييز ثلاثة أنواع رئيسية من التعلّم الآليّ:

< يتعلّم نظام الذكاء الاصطناعيّ التنبؤ بالقيمة اعتمادا على بيانات مدخلة خلال التعلّم الخاضع للإشراف. ويتطلّب ذلك وجود الأزواج مُدخل-قيمة المشروحة خلال التدريب. يمكن للنظام على سبيل المثال تعلّم التعرف على شيء موجود في صورة.

< يتعلّم نظام الذكاء الاصطناعيّ كيفية العثور على أوجه التشابه بين البيانات غير المشروحة، خلال التعلّم غير الخاضع للإشراف، من أجل تقسيمها على سبيل المثال إلى عدّة أقسام متجانسة. هكذا، يمكن للنظام التعرّف على مجتمعات مستخدمين وسائل التواصل الاجتماعيّ.

< يتعلّم نظام الذكاء الاصطناعيّ التّحكّم في بيئته خلال التعلّم المعزز، بطريقة ترفع لأقصى حدّ المكافأة التي يحصل عليها أثناء التدريب. وتتمثّل هذه التقنيّة التي تمكّن من خلالها أنظمة الذكاء الاصطناعيّ من التّغلب على بشر في لعبة Go أو لعبة الفيديو Dota 2.

التعلّم العميق

يمثّل التعلّم العميق أحد فروع التعلّم الآليّ الذي يستخدم شبكات عصبية اصطناعية على عدة مستويات، ويعتبر التقنيّة الكامنة وراء أحدث تطوّرات الذكاء الاصطناعيّ.

التّسمية المستدامة

تُشير التّسمية المستدامة إلى تطوّر للمجتمعات البشرية يتوافق مع قدرة النظم الطبيعية على توفير الموارد والخدمات اللازمة لهاته المجتمعات. وتتمثّل في تنمية اقتصادية واجتماعية تلبي الاحتياجات الحالية للأشخاص دون تعريض وجود الأجيال المقبلة للخطر.

الخوارزميات

تمثل الخوارزمية طريقة لحل المشكلات عبر سلسلة محدودة وغير مبهمه من العمليات. وبشكل أكثر تحديداً، في مجال الذكاء الاصطناعي تمثل الخوارزمية سلسلة العمليات المطبقة على بيانات الإدخال لتحقيق النتيجة المرجوة.

الذكاء الاصطناعي

يشير الذكاء الاصطناعي إلى سلسلة التقنيات التي تمكن الآلة من محاكاة الذكاء البشري، وذلك من أجل التعلم والتنبؤ واتخاذ القرارات وإدراك البيئة المحيطة بها، على سبيل المثال. ويطبق الذكاء الاصطناعي على البيانات الرقمية عندما يتعلق الأمر بالنظم المعلوماتية.

المشاعات الرقمية

تتمثل المشاعات الرقمية في التطبيقات أو البيانات التي تنتجها أحد المجتمعات. ويمكن مشاركتها بسهولة، كما أنها لا تتلف عند استخدامها، عكس السلع المادية. لذلك، وعلى عكس البرامج الاحتكارية، تمثل البرامج مفتوحة المصدر - والتي غالباً ما تنتج عن التعاون بين المبرمجين - مشاعاً رقمياً نظراً لأن شفرة برمجيتها مفتوحة، بما معناه أنها متاحة للجميع.

المصادقية

يعتبر نظام الذكاء الاصطناعي موثقاً به عندما يؤدي المهام التي صمم من أجلها بالطريقة المرجوة. وتمثل المصادقية احتمالية النجاح التي تتراوح بين ٥١٪ و ١٠٠٪، وهذا يعني أنها نسبة أعلى من أن تكون عشوائية. وكلما كان النظام موثقاً به، كلما كان بالإمكان التنبؤ بسلوكه.

المعرفة الرقمية

تشير المعرفة الرقمية للفرد إلى قدرته على التعامل مع المعلومات وفهمها ودمجها وإبلاغها وتقييمها وخلقها والوصول إليها بصفة آمنة وبشكل مناسب وذلك باستخدام الأدوات الرقمية والتقنيات الشبكية للمشاركة في الحياة الاقتصادية والاجتماعية.

المفهومية

يكون نظام الذكاء الاصطناعي مفهوماً عندما يكون فهم كيفية سير عمله، أي نموذج الحساب والعمليات التي تحدده ممكناً لأي إنسان لديه المعرفة اللازمة لذلك.

النشاط الرقمي

نعني بالنشاط الرقمي مجموعة الإجراءات التي يتخذها الفرد في بيئة رقمية، سواء كان ذلك على جهاز كمبيوتر أو هاتف أو أي شيء آخر متصل.

أنظمة التّقاط وأرشفة البيانات الشخصية (DAAS)

يمثل نظام التقاط وأرشفة البيانات الشخصية كل نظام حسابي يمكنه جمع البيانات وتسجيلها. وقد تُستخدم هذه البيانات لتدريب أنظمة الذكاء الاصطناعي، أو كمعايير لاتخاذ القرار.

انقطاع الاتصال الرقمي

يشير انقطاع الاتصال الرقمي إلى التوقف المؤقت أو الدائم للفرد عن ممارسته للنشاط الرقمي.

روبوت الدردشة (Chatbot)

روبوت الدردشة هو نظام يستخدم الذكاء الاصطناعي ويمكنه التّحاور مع مستخدميه بلغة طبيعية.

شبكات التنافس الإنتاجي (GAN)

في شبكات التنافس الإنتاجي تتنافس شبكتان متعارضتان لاستحداث صورة. ويمكن على سبيل المثال أن يتم استخدامهما لإنشاء صورة أو تسجيل أو فيديو ليبدو شبه حقيقي للإنسان.

فقاعة التصفية

يشير تعبير فقاعة التصفية إلى المعلومات «المصفاة» التي تصل الفرد عبر الإنترنت. وتوفر في الواقع مختلف الخدمات على غرار الشبكات الاجتماعية أو محركات البحث نتائج مخصصة لمستخدميها. ويمكن أن يؤدي هذا إلى عزل الأفراد (وضعهم داخل «فقاعات») بما أنهم يفقدون القدرة على الوصول إلى معلومات مشتركة.

مبررات القرار

يمكن تبرير قرار اتخذه نظام الذكاء الاصطناعي عندما تكون هناك أسباب هامة تحفّزه وعندما تكون هذه الأسباب قابلة للتبليغ بلغة طبيعية.

نظام الذكاء الاصطناعي

يمثل نظام الذكاء الاصطناعي أي نظام حسابي يستخدم خوارزميات الذكاء الاصطناعي، سواء كان ذلك برنامجاً أو جسماً متصلاً أو روبوتاً.

الاعتمادات

إن صياغة إعلان مونتريال للتنمية المسؤولة للذكاء الاصطناعي هو نتاج عمل فريق علمي متعدد الاختصاصات ومن جامعات مختلفة، يعتمد على مبدأ الاستشارة المواطنية والتشاور مع الخبراء والأطراف الفاعلة في تطوير الذكاء الاصطناعي.

Christophe Abrassart أستاذ مبرز في كلية التصميم والمدير المشارك لمختبر Lab Ville Prospective التابع لكلية التخطيط بجامعة مونتريال وعضو في مركز البحث في الأخلاقيات (CRÉ).

Yoshua Bengio أستاذ قار في قسم الإعلامية وبحوث العمليات (DIRO) التابع لجامعة مونتريال ومدير الشؤون العلمية في MILA و IVADO

Guillaume Chicoisne مدير البرامج العلمية في IVADO

Nathalie de Marcellis-Warin أستاذة قارة في جامعة بوليتكنيك مونتريال (Polytechnique Montréal) والرئيسة والمديرة التنفيذية لمركز البحوث في تحليل المنظمات المشترك بين الجامعات (CIRANO).

Marc-Antoine Dilhac أستاذ مبرز في قسم الفلسفة بجامعة مونتريال ومدير محور الأخلاقيات والسياسة في مركز البحث في الأخلاقيات (CRÉ) ومدير معهد الفلسفة والمواطنة والشباب (Institut Philosophie Citoyenneté Jeunesse) وهو يشغل سدة بحث كندية حول الأخلاقيات العامة والنظرية السياسية.

Sébastien Gambs أستاذ علوم الحاسوب بجامعة كيبيك في مونتريال وهو يشغل كرسي البحث الكندي في التحليل المحترم للحياة الخاصة وأخلاقيات البيانات الضخمة.

Vincent Gauthrais أستاذ قار في كلية الحقوق بجامعة مونتريال ومدير مركز أبحاث في القانون العام (CRDP).

Martin Gibert مستشار الأخلاقيات في IVADO وباحث في مركز البحث في الأخلاقيات (CRÉ).

Lyse Langlois أستاذة مبرزة ونائبة عميد كلية العلوم الاجتماعية بقسم العلاقات الصناعية التابع لجامعة لافال ومديرة معهد الأخلاقيات التطبيقية (IDÉA)، وباحثة في مجال العولمة والعمل في مركز البحوث المشتركة بين الجامعات (CRIMT).

François Laviolette أستاذ قار في قسم الإعلامية وهندسة البرمجيات بجامعة لافال، ومدير مركز أبحاث البيانات الكثيفة (CRDM).

Pascale Lehoux أستاذة قارة في مدرسة الصحة العامة بجامعة مونتريال (ESPUM)، متحصلة على سدة جامعة مونتريال للابتكار المسؤول في الصحة.

Jocelyn Maclure أستاذ قار في كلية الفلسفة بجامعة لافال ورئيس لجنة الأخلاقيات في العلوم والتكنولوجيا (CEST).

Marie Martel أستاذة مساعدة في مدرسة علوم المكتبات والمعلومات بجامعة مونتريال.

Joëlle Pineau أستاذة مبرزة في مدرسة علوم الحاسوب بجامعة ماكجيل ومديرة مختبر IA التابع لفيسبوك والمديرة المشاركة لمختبر التفكير والتعلم (Laboratoire Reasoning and Learning).

Peter Railton أستاذ جامعي متميز حاصل على أستاذية Gregory S. Kavka؛ John Stephenson Perrin؛ Arthur F. Thurnau أستاذ بقسم الفلسفة بجامعة ميشيغان وعضو الأكاديمية الأمريكية للفنون والعلوم.

Catherine Régis أستاذة مبرزة في كلية الحقوق بجامعة مونتريال ومتحصلة على سدة بحث كندية للثقافة التعاونية في القانون والسياسات الصحية وباحثة قارة في مركز البحوث في القانون العام (CRDP).

Christine Tappelet أستاذة قارة في قسم الفلسفة التابع لجامعة مونتريال ومديرة مركز البحث في الأخلاقيات (CRÉ) ومسؤولة عن مجموعة العمل المشتركة بين الجامعات والمعنية بالحياة الطبيعية.

Nathalie Voarino المنسقة العلمية للوثيقة وطالبة الدكتوراه في العلوم الطبية الحيوية إختصاص أخلاقيات علم البيولوجيا في جامعة مونتريال.

Université 
de Montréal



CENTRE DE RECHERCHE EN ETHIQUE



ICRA
Programme
IA et
société



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





</>

montrealdeclaration-responsibleai.com