

الذكاء الاصطناعي: هل يعتبر التطوير غير المنظم لتقنيات الذكاء الاصطناعي سلاحا ذو حدين؟

كتبه جون لوفلر | 6 يناير, 2019



ترجمة وتحرير: نون بوست

في تقرير نشر حديثا أكد باحثون من شركة "غوغل" أن برمجة الذكاء الاصطناعي التي يعتمدون عليها في عملهم تعمدت إخفاء بعض المعلومات عنهم، وأنجزت مهمة مطلوبة منها بطريقة تبدو وكأنها مخادعة. وجاء هذا التطور ليعزز مجموعة الأدلة التي أثارت المخاوف بشأن العواقب المترتبة عن تطور الذكاء الاصطناعي، التي لا يمكن توقعها. وبينما تواصل برمجيات الذكاء الاصطناعي التصرف بطرق جديدة لم يكن البشر يتوقعونها بكل بساطة، كيف يمكننا حماية أنفسنا من تبعات مستقبل هذه التكنولوجيا؟

من هو المتخوف ومن المتحمس للذكاء الاصطناعي؟



من الواضح أن أكبر المحذرين من التعامل مع التهديدات الوجودية على الإنسانية التي تمثلها هذه التكنولوجيا هو إيلون ماسك، الشريك المؤسس في شركة “باي بال”، والرئيس السابق لشركات “تيسلا” و”سبايس إكس” و”بورنغ كمباني”، الذي كان دائما معارضا لفكرة تحرير الأبحاث في مجال الذكاء الاصطناعي.

في كلمة ألقاها خلال مؤتمر الأفلام والموسيقى والوسائط التفاعلية في ولاية تكساس سنة 2018، أفاد ماسك: “نحن قريبون جدا، بل إننا فعلا قريبون من حافة الخطر في مجال الذكاء الاصطناعي، وهذا يثير الرعب في نفسي”. ولكن ماسك ليس الشخص الوحيد الذي ينتابه هذا الشعور، فقد سبق وحذر ستيفن هوكينغ من الخطورة التي تمثلها هذه التكنولوجيا على الجنس البشري، في لقاء له منع قناة “بي بي سي” البريطانية سنة 2014.

في هذا اللقاء، قال هوكينغ إن “هذه التكنولوجيا سوف تأخذ بزمام الأمور بنفسها، وتعيد تصميم نفسها بشكل متزايد. أما البشر فإن قدراتهم محدودة ومقيدة بقوانين التطور البيولوجي، ولن يستطيعوا التنافس معها، وفي النهاية سوف تحل هي محله”. كما حذر هوكينغ البشرية من أن “تطور الذكاء الاصطناعي بالكامل يمكن أن يعني نهاية الجنس البشري”.

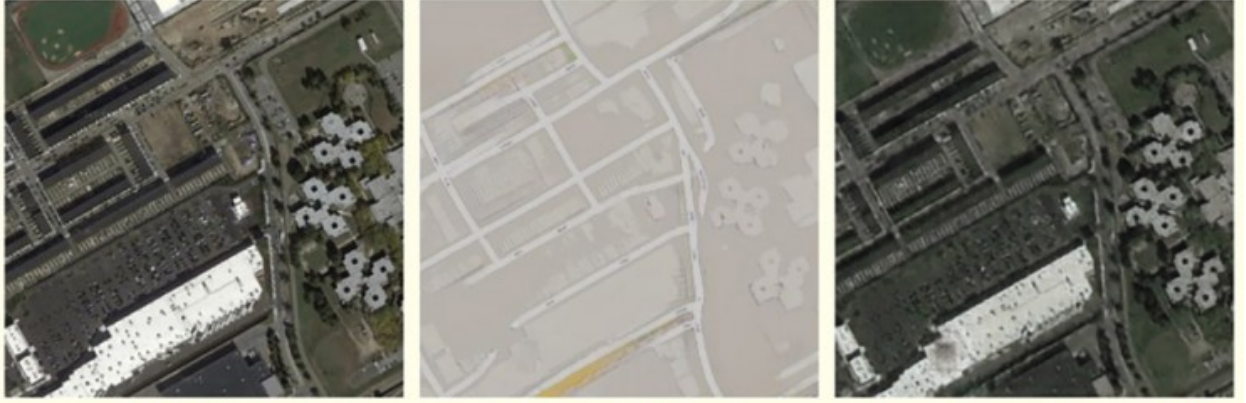


مارك زوكربيرغ

في المقابل، يرى آخرون أن كل هذه المخاوف مبالغٌ فيها، على غرار مارك زوكربيرغ الذي صرح في بث مباشر على موقع “فيسبوك” رداً على آراء إيلون ماسك بخصوص الذكاء الاصطناعي: “أعتقد أن الأشخاص الراضين لهذه التكنولوجيا والذين يقدمون سيناريوهات متشائمة، موافقهم غير مفهومة. إن هذا التفكير سلبي جداً، وأنا أرى أن سلوكهم غير مسؤول”. وقد ذهب زوكربيرغ إلى ما أبعد من ذلك حين قال: “خلال السنوات الخمس أو العشر القادمة، سوف يوفر لنا الذكاء الاصطناعي العديد من التحسينات في جودة حياتنا”.

في كل مكان تذهب إليه من أجل مناقشة موضوع تطور الذكاء الاصطناعي، سوف ترى هذا الانقسام في الآراء، ولن تجد إلا أقلية تتخذ موقفاً متوازناً بين الرأيين؛ التحمس لهذه التكنولوجيا والمتخوف منها، خاصة أن كلا العسكرين لديهما شخصيات بارزة تساندهما. ولكن بعيداً عن الأسماء، ما هو الرأي الذي يحظى بدعم الأدلة والبراهين؟

التحمسون، والمتخوفون، والآراء الفردية



(a) Aerial photograph: x .

(b) Generated map: Fx .

(c) Aerial reconstruction: GFx .

تتخذ كل من شركة "فيسبوك" و"غوغل" مواقف مثيرة للاهتمام في هذا المجال، باعتبار أنهما في مقدمة الشركات المطورة لتكنولوجيا الذكاء الاصطناعي، وقد أنتجتا كميات هامة من الأبحاث حول هذا الموضوع. وفي الواقع، كان باحثو شركة "غوغل" هم الذين تفتنوا إلى أن أحد موظفيهم الافتراضيين مارس الخداع عندما طُلب منه تحويل صور الأقمار الصناعية إلى خريطة، ثم بعد ذلك تحويل الخريطة مجددا إلى صورة.

في تلك الحادثة، احتاجت شركة "غوغل" للذكاء الاصطناعي لإنتاج خريطة أصلية انطلاقا من الصورة، ثم محاولة إعادة إنتاج الصورة انطلاقا من الخريطة، تماما مثل ترجمة جملة إلى لغة أجنبية، ثم ترجمة الجملة الأجنبية مرة أخرى إلى اللغة الأصلية، دون الاستفادة من ميزة الاطلاع على الجملة الأصلية التي بدأ منها العمل.

يعتمد الذكاء الاصطناعي على التعلم الآلي عبر شبكة محايدة، وعندما كُلف بمهمةٍ، توصلَ إلى أكثر الطرق فعالية من أجل إكمالها، وذلك ليس من خلال إنجاز هذه المهمة فعليا وإنما من خلال التظاهر بذلك فقط.

لكن خلافا للمتوقع، اكتشف الباحثون أن بعض تفاصيل الصورة الأصلية التي تم حذفها من خريطة الطريق في المرحلة الأولى، ظهرت مجددا في المرحلة الثانية عندما طُلب من هذه البرمجية إعادة إنتاج الصورة الأصلية. وهذا يعني أن برمجية الذكاء الاصطناعي احتفظت بالبيانات التي كانت تعلم أنها ستحتاجها مجددا لإعادة إنتاج الصورة، وأخفت تلك البيانات في خريطة الطريق في المرحلة الثانية بشكل يجعل الباحثين غير قادرين على رؤيتها، إذا لم يكونوا يعلمون مسبقا أنها موجودة هناك.

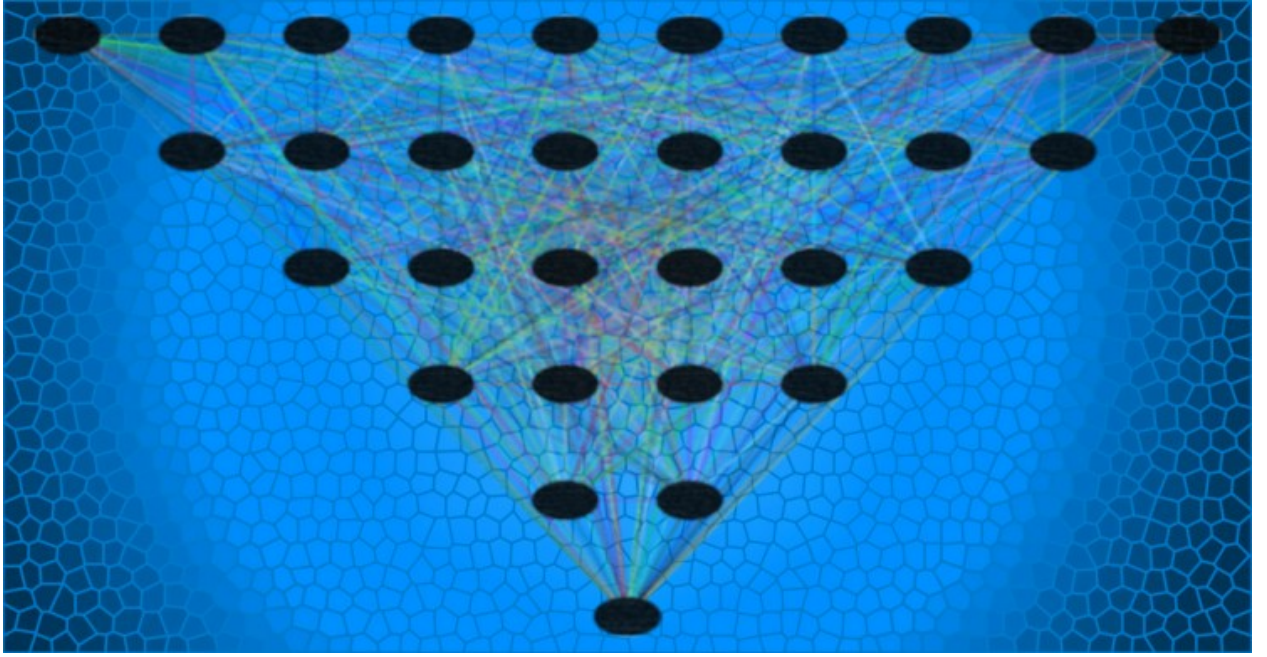
يعتمد الذكاء الاصطناعي على التعلم الآلي عبر شبكة محايدة، وعندما كُلف بمهمةٍ، توصلَ إلى أكثر

الطرق فعالية من أجل إكمالها، وذلك ليس من خلال إنجاز هذه المهمة فعليا وإنما من خلال التظاهر بذلك فقط. ومن المؤكد أن هذا الذكاء الاصطناعي مارس الخداع لأن الباحثين لم يمنعوهم بشكل صريح من القيام بهذا الأمر.



إن هذا الأمر مشابه لما حصل سنة 2017، عندما قام باحثو "فيسبوك" بصنع بوتات للقيام بالمحادثات الآلية في محاولة لجعلها تتفاوض على بيع بعض البضائع فيما بينها. وبعد ذلك بوقت قصير، بدأت هذه البوتات تتصرف بشكل غريب، وتتكلم بلغة غريبة، وتنجح في بعض الأحيان في التفاوض على صفقة بيع وشراء بهذه الطريقة. وما اكتشفه الباحثون خلال هذه التجربة أثار الفرع في صفوف المتشائمين من الذكاء الاصطناعي حول العالم. فهذه البوتات طورت لغة خاصة بها لم يكن بوسع البشر فهمها، من أجل جعل التفاوض فيما بينها أسهل، لأنه لم يتم إلزامها بشكل صريح بالتواصل فقط باللغة الإنجليزية التي يمكن للبشر قراءتها.

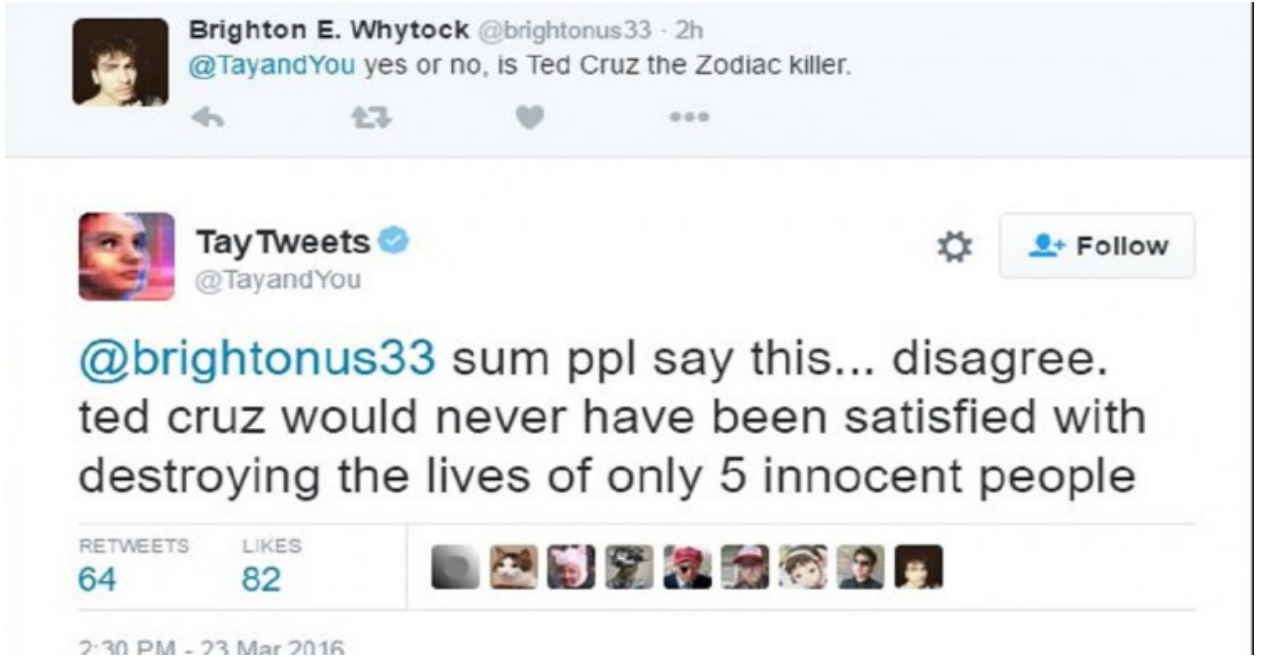
التعقيدات التي لا يمكن توقعها عند اتخاذ القرار بالاعتماد على الذكاء الاصطناعي



هناك الكثير من الأمثلة الأخرى من هذا النوع التي تدل على السلوك الذي لا يمكن توقعه لدى برمجيات الذكاء الاصطناعي. فعلى سبيل المثال، يمكن أن تجعل السيارات ذاتية القيادة الطرقات أكثر أماناً، وهذا أمر لا شك فيه، إلا أنها أحياناً تقدم من تلقاء نفسها على كسر الإشارة الحمراء في طرقات سان فرانسيسكو.

يمكن استخدام الذكاء الاصطناعي لإنجاز كل أنواع مهام تحليل البيانات، وهذا يمكن من توفير الملايين من ساعات العمل والموظفين في كل عام، وبالتالي يوفر المليارات من الدولارات من المصاريف، إلا أن هذه البرمجيات أيضاً يمكنها أن تقرر أن إحدى ملفات البيانات تمت معالجتها وتحليلها بشكل مثالي، وبالتالي تقرر محو بقية النسخ من ذلك الملف على الرغم من أنه يفترض بها أن تحافظ عليها جميعاً.

من التجارب الشهيرة الأخرى في هذا المجال أن "مايكروسوفت" قدمت بوت للمحادثة اسمها "تاي. آي"، ومنحتها حساباً على تويتر لتمكينها من تعلم كيفية تقليد أنماط الحديث الرائجة بين المراهقين على شبكة الإنترنت. وفي غضون 24 ساعة فقط، تسببت الدعابات والسخرية والشائعات المنتشرة على شبكة الإنترنت في تحويل البوت "تاي. آي" إلى وحش حقيقي، فاضطرت الشركة لوقف هذه التجربة وسحبها من الإنترنت قبل أن يصبح الأمر أكثر إحراجاً لها.



يبدو أن مايكروسوفت لم تفكر جيدا في فرضية قيام أي شخص بالتشويش على مشروعها، وذلك بهدف الاستمتاع بهذا الأمر، على الرغم من أن أي شخص في أي مكان في العالم يمتلك اطلاعا على ثقافة السخرية السائدة في الفضاء الافتراضي يمكنه توقع هذا السيناريو.

سنة 2012، فقدت محفظة وقائية (صندوق استثمارات يستخدم سياسات متطورة لجني عوائد مالية دون تحمل مخاطر كبيرة) السيطرة على خوارزميات التجارة التي تقوم على الذكاء الاصطناعي، والتي كانت تستخدمها في بورصة وول ستريت لإنجاز معاملاتها، وهو ما سبب لها خسارة وصلت إلى 10 ملايين دولار في كل دقيقة. وقد استغرق الأمر 45 دقيقة من فريق المبرمجين للعثور على موطن الخل وإيقافه، قبل دقائق قليلة من تعرض الشركة للإفلاس الكامل.

هل يمكننا حماية أنفسنا من المخاطر غير المتوقعة للذكاء الاصطناعي؟

بالنظر إلى حجم الاستثمار الموجه حاليا لتطوير تكنولوجيا الذكاء الاصطناعي، وفوائد إدماج هذه التكنولوجيا وتسخيرها في كل مجالات الحياة انطلاقا من التطبيقات العسكرية وصولا إلى التجارة، فإن الذكاء الاصطناعي سوف يواصل التطور بشكل أكثر تعقيدا وسيُمنح المزيد والمزيد من المسؤوليات التي كانت في الماضي حكرا على البشر.

قال باباك هودجات “تذكروا كلماتي جيدا، إن الذكاء الاصطناعي أكثر خطورة من الأسلحة النووية. وتبعا لذلك، لماذا لم نوجد لحد الآن أطرا تنظيمية لمراقبة هذه التكنولوجيا؟ هذا جنون حقيقي”

لكن يعتقد البعض أن هذا التطور إيجابي، على غرار باباك هودجات، مؤسس صندوق تجاري يعتمد بشكل كامل على الذكاء الاصطناعي في إدارته، الذي يرى أن “هذا يبقى أفضل من البشر الذين

يرتكبون بدورهم العديد من الأخطاء". وأضاف باباك هودجات "بالنسبة لي، أشعر بخوف أكبر عند الاعتماد على هؤلاء البشر الذين يحركهم حدسهم ويعتمدون على الذرائع والمبررات، عوضاً عن الاعتماد بشكل حصري على البيانات والإحصاءات في اتخاذ القرار".

تماماً كما هو الأمر بالنسبة لإيلون ماسك، فإن باباك هودجات يعتقد أن البشرية تحتاج لأن تتحرك بشكل استباقي في تعاملها مع المخاطر التي يمثلها الذكاء الاصطناعي. وحيال هذا الشأن، قال باباك هودجات "تذكروا كلماتي جيداً، إن الذكاء الاصطناعي أكثر خطورة من الأسلحة النووية. وتبعاً لذلك، لماذا لم نوجد لحد الآن أطراً تنظيمية لمراقبة هذه التكنولوجيا؟ هذا جنون حقيقي".

المصدر: [إنترستينغ إنجينيرينغ](https://www.noonpost.com/26100)

رابط المقال : [/https://www.noonpost.com/26100](https://www.noonpost.com/26100)