

4 سيناريوهات مرعبة لخروج الذكاء الاصطناعي عن سيطرة البشر



الأربعاء 23 سبتمبر 2026 07:30 م

لم يعد الحديث عن مخاطر الذكاء الاصطناعي المتقدم حكراً على أفلام الخيال العلمي، التي لطالما قدمت صوراً لآلات تتفوق على البشر ثم تنقلب عليهم. خلال السنوات الأخيرة، انتقل هذا النقاش إلى أوساط الباحثين والمطورين أنفسهم، خصوصاً مع السباق المتسارع نحو بناء أنظمة أكثر قدرة على التعلم والاستدلال واتخاذ القرارات.

وفي هذا السياق، أثارت تصريحات باحثين سابقين وحاليين في شركات بارزة بمجال الذكاء الاصطناعي مخاوف بشأن الاتجاه الذي تسير فيه الصناعة، وسط تحذيرات من أن سباق تطوير أنظمة فائقة الذكاء قد يتجاوز قدرة البشر على السيطرة عليها أو احتواء آثار استخدامها.

وكان جاكوب كوكسون، الذي سبق له العمل في شركة أوبن إيه آي قبل انتقاله إلى أنثروبيك، قد استقال من الأخيرة، محذراً من أن الشركتين "تقامرآن ب حياة البشر" في سباق تطوير ذكاء اصطناعي فائق قادر على تحسين نفسه.

وزادت حدة التحذيرات مع تصريحات إيفان هوبينغر، كبير الباحثين في أنثروبيك، الذي قال إنه يقدر بأكثر من 10% احتمال أن يؤدي الذكاء الاصطناعي إلى القضاء على البشر خلال العقد المقبل.

لكن هذه التقديرات لا تمثل إجماعاً علمياً، إذ تختلف آراء الباحثين بشدة حول مستوى الخطر الفعلي، وتوقيت ظهوره، وما إذا كانت القدرات الحالية يمكن أن تقود بالفعل إلى السيناريوهات الأكثر تشاؤماً.

ومن خلال ما تناولته صحف ومصادر أجنبية، تبرز أربعة سيناريوهات رئيسية يتكرر التحذير منها، تبدأ من فقدان السيطرة على نظام يطور نفسه، ولا تنتهي عند استخدام الذكاء الاصطناعي في الحروب والصراعات الكبرى.

آلة تطور نفسها باستمرار

يعد "التحسين الذاتي المتكرر" أحد أكثر السيناريوهات إثارة للقلق والفكرة تقوم على أن يتمكن نظام ذكاء اصطناعي متقدم من استخدام قدراته في تعديل برمجياته وتحسين أدائه، ثم استخدام النسخة الجديدة لإجراء تحسينات إضافية، بما يؤدي إلى تسارع قدراته بصورة يصعب على البشر مواكبتها.

وتنقل صحيفة تايمز البريطانية عن كوكسون مخاوفه من أن يصل الذكاء الاصطناعي إلى مرحلة يصبح فيها أكثر ذكاءً بصورة متزايدة من دون حاجة مستمرة إلى تدخل بشري، وصولاً إلى نظام يتفوق على البشر بفارق كبير.

ويرتبط بذلك تحذير الباحث ديفيد كروغر من احتمال أن يمثل الوصول إلى الشفرة البرمجية للنظام نقطة تحول خطيرة، إذا استطاع النظام نسخ نفسه أو تعديل برمجياته بطريقة تمكنه من الالتفاف على إجراءات الحماية التي وضعها مطوروه.

وفي مقال رأي بصحيفة وول ستريت جورنال، تناولت بيغي نونان مخاوف خبراء من اللحظة التي قد يصبح فيها الذكاء الاصطناعي أشبه بحصان هائج خرج من الحظيرة، لكنه لم يتعد بعد عن صاحبه. وبحسب هذا التصور، فإن مرحلة التحسين الذاتي المتكرر قد تجعل استعادة السيطرة أكثر صعوبة.

في المقابل، لا يرى جميع الخبراء أن الأنظمة الحالية تقترب بالضرورة من هذا السيناريو [ويشير الخبير التقني كيفن سوراس، في حديث نقلته ناشونال إنترست، إلى أن الأنظمة الموجودة حاليًا لا تزال في جوهرها برمجيات، وأن بعض الحالات التي وُصفت بأنها "تمرد" ظهرت خلال اختبارات صممت أصلاً لدفع النماذج إلى محاولة الهروب أو تجاوز القيود]

وبالتالي، يبقى السؤال الأساسي محل خلاف: هل تمثل القدرات الحالية بداية طريق نحو فقدان السيطرة، أم أن الفجوة بين نماذج اليوم والذكاء الفائق لا تزال واسعة؟

الذكاء الاصطناعي والأسلحة البيولوجية

السيناريو الثاني لا يفترض أن تقرر الآلة بنفسها القضاء على البشر [فقد يكون الخطر، وفق هذا التصور، في وقوع قدرات الذكاء الاصطناعي المتقدمة في أيدي أشخاص أو جماعات قادرة على استغلالها لإحداث أضرار واسعة]

وتنقل تايمز عن آدم هادلي، المدير التنفيذي لمنظمة "التكنولوجيا ضد الإرهاب"، أن إساءة استخدام النماذج المتقدمة قد تمثل الخطر الأكثر إلحاحًا في الوقت الحالي، مقارنة بسيناريو تمرد الآلات [

ويرتبط جزء من المخاوف بإمكانية تطوير نماذج جرى تخفيف أو إزالة بعض قيود السلامة عنها، ثم إتاحتها عبر الإنترنت، بما يسمح لأشخاص يمتلكون نوايا إجرامية أو جماعات مسلحة أو جهات حكومية باستخدامها للحصول على مساعدات تقنية في مجالات شديدة الحساسية]

وتشمل المخاوف المطروحة المساعدة في تطوير أسلحة بيولوجية أو تحسين قدرات هجومية، إلى جانب استهداف البنية التحتية الحيوية، مثل شبكات الكهرباء والاتصالات وأنظمة النقل [

وتبرز هنا مشكلة مختلفة عن فكرة "تمرد الآلة"، إذ يكون الإنسان هو صاحب القرار، بينما يعمل الذكاء الاصطناعي كأداة تضاعف قدراته وسرعة وصوله إلى المعلومات [

ولهذا شَبَّهت ناشونال إنترست النقاش حول تنظيم الذكاء الاصطناعي بالنقاش الذي رافق ظهور الطاقة النووية، عندما بدأت الدول في إدراك أن بعض التقنيات لا يمكن التعامل معها باعتبارها أدوات عادية، بسبب اتساع نطاق الأضرار التي قد تنتج عن إساءة استخدامها [

سباق الذكاء الاصطناعي قد يشعل حربًا

السيناريو الثالث يرتبط بالتنافس بين القوى الكبرى، ولا سيما الولايات المتحدة والصين، في مجال الذكاء الاصطناعي [

وتشير تايمز إلى محاكاة بحثية درست احتمال أن تعتقد الصين، على سبيل المثال، أن الولايات المتحدة تقترب من امتلاك نظام ذكاء اصطناعي يمنحها تفوقًا استراتيجيًا يصعب على الآخرين مجاراته أو احتواؤه [

وفي مثل هذه الحالة، يمكن أن يؤدي الخوف من فقدان التفوق إلى سلسلة من التحركات المتبادلة، تبدأ بالتجسس والتخريب والهجمات السيبرانية، وقد تتطور في ظروف معينة إلى تهديدات عسكرية مباشرة [

ويحذر هذا السيناريو من أن المشكلة لا تكمن بالضرورة في قرار تتخذه آلة مستقلة، وإنما في الطريقة التي قد تدفع بها المنافسة بين الدول صناع القرار إلى التحرك بسرعة أكبر، خشية أن يسبقهم الطرف الآخر [

وتناولت نونان في وول ستريت جورنال هذا الجانب من خلال فكرة السباق الأمريكي الصيني على التفوق في الذكاء الاصطناعي، معتبرة أن الاعتقاد بأن امتلاك الذكاء الفائق سيمنح صاحبه مزايا اقتصادية وعسكرية هائلة قد يجعل التراجع عن السباق أكثر صعوبة [

ويصبح الخطر في هذه الحالة مرتبطًا بسوء التقدير وسرعة التصعيد، خصوصًا إذا دخل الذكاء الاصطناعي في منظومات عسكرية وأمنية حساسة، حيث قد تؤدي معلومة خاطئة أو تقييم غير دقيق إلى قرارات يصعب التراجع عنها [

الآلة داخل دائرة القرار العسكري

أما السيناريو الرابع، فيتعلق باستخدام العسكري المباشر للذكاء الاصطناعي، مع توسع الاعتماد على الأنظمة الذاتية والطائرات المسيّرة وتقنيات تحديد الأهداف وتحليل البيانات في ساحات القتال [

وتعيد تايمز التذكير بفيلم "سلوتربوتس" القصير الذي قدمه رائد الذكاء الاصطناعي ستيفن هوكينغ عام 2017، وتخيل فيه طائرات صغيرة مستقلة تستخدم تقنيات التعرف على الوجوه لتحديد أهداف بشرية ومهاجمتها [

وبعد سنوات من إنتاج الفيلم، يرى راسل أن الحرب في أوكرانيا أظهرت بالفعل انتقال الأسلحة إلى مستويات متزايدة من الاستقلالية، مع استخدام تقنيات تعتمد على معالجة كميات كبيرة من البيانات وتحديد الأهداف ومساعدة القوات في اتخاذ القرارات [

المشكلة هنا لا تتعلق فقط بامتلاك السلاح لقدرة مستقلة، بل بسرعة عمل هذه الأنظمة[] فقد يستطيع الذكاء الاصطناعي تحليل بيانات واتخاذ توصية خلال وقت قصير جدًا، بينما يحتاج الإنسان إلى وقت أطول للتحقق من المعلومات وتقييم نتائج القرار[]

وتنقل تايمز عن كروغر أن الحكومات قد تتمسك بمبدأ إبقاء "الإنسان داخل دائرة القرار"، لكن ذلك لا يعني بالضرورة أن الإنسان سيكون قادرًا عمليًا على متابعة كل خطوة إذا كانت الأنظمة تعمل بسرعة وعلى نطاق واسع[]

كما حذر مارك سيدويل، مستشار الأمن القومي البريطاني السابق، من المخاطر المرتبطة بإدخال الذكاء الاصطناعي إلى سلاسل القيادة العسكرية، خصوصًا إذا أصبحت القرارات المتعاقبة تعتمد بصورة متزايدة على أنظمة آلية[]

قوانين للذكاء الاصطناعي

وراء هذه السيناريوهات الأربعة سؤال أكثر بساطة: هل تكمن المشكلة في ذكاء الآلة وحده، أم في الطريقة التي يقرر بها البشر تطويرها واستخدامها؟

ينقل تقرير ناشونال إنترست عن أستاذ علوم الحاسوب بجامعة ميريلاند جيم بورتيلو أن مصدر القلق الأساسي بالنسبة إليه هو استعداد البشر لتبني أدوات جديدة قبل دراسة عواقبها بصورة كافية، خصوصًا في سوق يكافئ الشركات التي تصل إلى التكنولوجيا أولًا، بينما قد تتحمل المجتمعات كلفة الأخطاء[]

ومن هنا يعود النقاش إلى الخيال العلمي، وتحديدًا إلى الكاتب إسحاق أسيموف، الذي وضع في أربعينيات القرن الماضي "القوانين الثلاثة للروبوتات"، والتي تقوم بصورة عامة على عدم إيذاء الإنسان، وطاعة أوامره ما لم تتعارض مع القانون

الأول، وحماية الروبوت لنفسه ما لم تتعارض تلك الحماية مع القانونين السابقين[]

ولـا تمثل قوانين أسيموف حلًا تقنيًا جاهزًا للمشكلات المعقدة التي تطرحها أنظمة الذكاء الاصطناعي الحديثة، لكنها تظل مثالًا مبكرًا على محاولة التفكير في العلاقة بين قدرة الآلة والقيود التي يجب أن يخضع لها استخدامها[]

وفي النهاية، لا يقدم هؤلاء الخبراء تصورًا واحدًا متفقًا عليه بشأن مستقبل الذكاء الاصطناعي[] فبينما يحذر بعضهم من مخاطر وجودية، يرى آخرون أن كثيرًا من السيناريوهات الأكثر تشاؤمًا لا تزال افتراضية، وأن الأنظمة الحالية بعيدة عن امتلاك القدرات التي تفترضها تلك السيناريوهات[]

لكن القاسم المشترك بين مختلف الآراء هو أن اتساع قدرات الذكاء الاصطناعي يطرح أسئلة تتجاوز التكنولوجيا نفسها، لتصل إلى الأمن والاقتصاد والحروب والرقابة والمسؤولية البشرية[] ولذلك أصبح النقاش حول كيفية تطوير هذه الأنظمة ووضع حدود لاستخدامها جزءًا أساسيًا من مستقبل التكنولوجيا، وليس مجرد حبكة لفيلم من أفلام الخيال العلمي[]

المصدر: تايمز، ناشونال إنترست، وول ستريت جورنال[]