

علماء آليون

دكتور صلاح عثمان

2024-09-17

العلم ليس مجرد مجموعة من الوقائع، لكنه عملية كشف مُمنهجة مدفوعة بالدهشة؛ والكشف العلمي ثمرة ناضجة لبحثٍ علمي ناجح. وثمار البحث متنوعة، تشمل الأحداث أو العمليات أو الأسباب، بالإضافة إلى الفروض والنظريات وخصائصها (القوة التفسيرية والقدرة على التنبؤ والديناميكية). لذا تُركز أغلب المناقشات الفلسفية حول الكشف العلمية على كيفية توليد فروض جديدة تناسب أو تُفسر مجموعة من المعطيات، أو تسمح باستنباط نتائج قابلة للاختبار.

بعبارة أخرى، نستطيع القول إن الكشف العلمي عملية معقدة وإنسانية للغاية، بل هو أحد أكثر الأنشطة البشرية تعقيدًا، حيث يجب على العلماء أولاً فهم المعرفة السائدة والوقوف على ماهية الفجوة القائمة – إن كانت ثمة فجوة – بينها وبين المعطيات الجديدة؛ ويجب عليهم ثانيًا صياغة سؤال بحثي، وتصميم تجربة وإجرائها سعيًا للحصول على إجابة مُشبعة؛ ويجب عليهم ثالثًا تحليل وتفسير نتائج التجربة، والتي قد تثير سؤالاً بحثيًا آخر؛ فهل بإمكاننا أتمتة عملية معقدة كهذه؟ وإذا كان ذلك مُمكنًا، فما ملامح هذه الأتمتة؟ وما تأثيرها على البحث العلمي؟ وما القيود التي تحكمها أو التحديات التي تُوجهها؟ هذا ما نسعى إلى تفصيله ومناقشته في هذا المقال.

قُبيل منتصف أغسطس 2024، أعلنت شركة "ساكانا للذكاء الاصطناعي" Sakana AI (وهي شركة ناشئة مقرها طوكيو باليابان) عن تدشين ما يُسمى «عالم ذكاء اصطناعي»، وهو نظام ذكاء اصطناعي يزعم مُطوروه أن بإمكانه إنجاز كشوفٍ علمية جديدة في مجال التعلم الآلي بطريقة مُؤتمتة بالكامل. وباستخدام نماذج اللغة التوليدية الكبيرة Generative Large Language Models (LLMs)، مثل تلك التي يعمل بها «تشات جي بي تي» وغيره من روبوتات الدردشة التي تعمل بالذكاء الاصطناعي، يمكن للنظام فحص مجموعة مختلفة من الأفكار، ثم اختيار فكرة واحدة، وترميز خوارزميات جديدة، وإجراء التجارب اللازمة، والتوصل إلى نتائج، وكتابة ورقة بحثية كاملة مُذيلة بالمراجع تلخص التجربة ونتائجها. ليس ذلك فحسب، بل ومراجعتها آليًا من قبل الأقران لتقييمها بدقة شبه بشرية، وكتابة الملاحظات، وتحسين النتائج بشكل

أكبر. كما تتكرر عملية الكشف العلمي الآلي لتطوير الأفكار بطريقة مفتوحة وإضافة ما يتم التوصل إليه إلى أرشيف متنامٍ من المعرفة، وبالتالي تقليد المجتمع العلمي البشري. وتزعم شركة «ساكانا» أن أداة الذكاء الاصطناعي الجديدة يمكنها أتمتة دورة حياة كاملة لتجربة علمية بتكلفة خمسة عشر دولاراً أمريكياً فقط لكل ورقة – أي بأقل من تكلفة الغداء لعالم بشري!

هذا بلا شك ادعاء كبير يبدو عصياً على التحقق، على الأقل حتى الآن. لكن ماذا لو كان من الممكن تحقيقه؟

هل وجود جيش من العلماء الآليين الذين يستطيعون إنتاج أوراق بحثية كاملة بسرعة تتخطى إمكانيات البشر يُعد خبراً جيداً بالفعل للعلم؟!

جاء إعلان "ساكانا" عن هذه الخطوة الرائدة في ورقة بحثية أعدها فريقها البحثي بالتعاون مع مختبر فورستر لأبحاث الذكاء الاصطناعي في جامعة أكسفورد، وعالمي الحاسوب "جيف كلون" و "كونج لو" من جامعة كولومبيا البريطانية، تحت عنوان عالم الذكاء الاصطناعي: نحو كشف علمي مفتوح ومؤتمت بالكامل، نُشرت في الثاني عشر من أغسطس وبينما صرّحت «ساكانا» بأنها لا ترى أن دور العلماء البشريين يتضاءل، إلا أنها تؤكد على أن رؤيتها لنظام بيئي علمي مدفوع بالكامل بالذكاء الاصطناعي سيكون له آثار كبيرة على العلم.

لا شك أن أحد التحديات الكبرى التي تواجه الذكاء الاصطناعي هو تطوير وكلاء قادرين على إجراء البحوث العلمية واكتشاف المعرفة الجديدة. وفي حين تم بالفعل استخدام نماذج الأساس لمساعدة العلماء البشريين، على سبيل المثال، لتبادل الأفكار أو كتابة التعليمات البرمجية، إلا أنها لا تزال تتطلب إشرافاً يدوياً مكثفاً أو مقيداً بشدة بمهمة محددة. ومن المعروف أن نماذج الأساس هي نماذج للتعليم الآلي أو التعلم العميق يتم تدريبها على بيانات ضخمة بحيث يمكن تطبيقها عبر مجموعة واسعة من حالات الاستخدام، وتشغيل نماذج اللغات التوليدية الكبيرة مثل «تشات جي بي تي».

لعمد خلت، وبعد كل تقدم كبير في مجال الذكاء الاصطناعي، كان من الشائع أن يسخر باحثوا الذكاء الاصطناعي فيما بينهم قائلين "الآن كل ما نحتاج إلى فعله هو معرفة كيفية جعل الذكاء الاصطناعي يكتب لنا الأوراق البحثية". لكن ظهور "عالم الذكاء الاصطناعي" الجديد يؤكد أن هذه الفكرة انتقلت من مجرد نكتة خيالية غير واقعية، بدت للباحثين مُضحكة، إلى شيء ممكن حالياً، كما أن تكلفة إنتاج الورقة البحثية، والوعود الذي يُظهره النظام حتى الآن يوضحان الإمكانيات الهائلة للذكاء الاصطناعي، وكيف يأخذنا إلى عالم حيث يمكن إطلاق العنان للإبداع والابتكار بأسعار معقولة لا نهاية لها لحل أكثر مشاكل العالم تحدياً.

أولاً: عالم الذكاء الاصطناعي (نظرة عامة)

"عالم الذكاء الاصطناعي" بمثابة خط إنتاج آلي للبيانات لتوليد أوراق بحثية كاملة من الألف إلى الياء. وبفضل التطورات الأخيرة في نماذج الأساس، وفي معية الاتجاهات البحثية واسعة النطاق التي تحتويها قواعد البيانات، مثل قاعدة البيانات مفتوحة المصدر للأبحاث السابقة على منصة «جيت هاب» (التابعة لشركة ميكروسوفت)، يقوم «عالم الذكاء الاصطناعي» أولاً بالعصف الذهني الآلي لمجموعة من الأفكار، ثم يُقيم مدى حداثتها، وبعدها يقوم بتحرير قاعدة بيانات مدعومة بالتطورات الحديثة في توليد التعليمات البرمجية الآلية لتنفيذ الخوارزميات الجديدة، ثم يقوم بإجراء تجارب لجمع النتائج المكونة من بيانات رقمية وملخصات مرئية، ثم يقوم بصياغة تقرير علمي يشرح النتائج ويضعها في سياقها، وأخيراً، يقوم «عالم الذكاء الاصطناعي» بإنشاء مراجعة أقران آلية بناءً على معايير مؤتمرات التعلم الآلي من الدرجة الأولى، حيث تساعد هذه المراجعة في تحسين المشروع الحالي وتزويد الأجيال القادمة بأفكار مفتوحة (انظر الشكل). وعلى هذا يجتاز «عالم الذكاء الاصطناعي» أربع مراحل نوعية أساسية.

توليد الفكرة: يقوم «عالم الذكاء الاصطناعي» أولاً بفحص الأفكار حول مجموعة متنوعة من اتجاهات البحث الجديدة. ويتم توفير قالب رمزي مبدئي لموضوع مُلح نرغب في أن يستكشفه بشكل أكبر، ثم يغدو حراً في استكشاف أي اتجاه بحثي ممكن. يتضمن القالب أيضاً مجلد لاتكس LaTeX (وهو برنامج مجاني لإعداد وكتابة الأبحاث والكتب والتقارير والرسائل العلمية المنشورة عالمياً، موجهٌ لذوي التخصصات العلمية كالرياضيات والفيزياء والحاسب والهندسة نظراً لدعمه الشامل لكتابة المعادلات الرياضية وتنسيقها)، ويحتوي المجلد على ملفات الأنماط (الملفات التي لا تهدف إلى إنتاج مخرجات، ولكنها تحدد بارامترات تخطيط المستند والأوامر والبيئات وما إلى ذلك)، وعناوين الأقسام الأساسية والفرعية لكتابة الأوراق البحثية. كذلك نسمح له بالبحث في «سيمانتيك سكولار» (الباحث الدلالي) Semantic Scholar للتأكد من أن فكرته جديدة.

تنفيذ التجارب: في المرحلة الثانية يقوم «عالم الذكاء الاصطناعي» بتنفيذ التجارب المقترحة، ثم الحصول على مخططات رقمية وبيانية وإنتاجها، ثم يقوم بعمل ملاحظة تصف ما يحتويه كل مخطط، مما يتيح للأشكال المحفوظة والملاحظات التجريبية توفير جميع المعلومات المطلوبة لكتابة الورقة. وكتابة الورقة البحثية وخطوة قبل أخيرة، يقوم «عالم الذكاء الاصطناعي» بكتابة ورقة مُركزة وغنية بالمعلومات عن تقدمه في المسار البحثي وفقاً للمعايير القياسية، ويستخدم «سيمانتيك سكولار» للعثور بشكل مستقل على أوراق بحثية ذات صلة للاستشهاد بها.

المراجعة الآلية للورقة البحثية: أحد الجوانب الرئيسية لهذا العمل هو تطوير مُراجع آلي يعمل بنظام النموذج اللغوي الكبير، وقادر على تقييم الأوراق البحثية المؤلّدة بدقة شبه بشرية، ما يسمح بتحسين المشروع أو ترك ملاحظات للأجيال القادمة من أجل ابتكار أفكار مفتوحة للبحث. تتيح هذه الخطوة ما يمكن أن نسميه حلقة ملاحظات مستمرة، من شأنها تحسين النتائج بشكلٍ متكرر. وفي هذا الصدد، تُسلط «ساكانا» الضوء على نماذج للأوراق البحثية التي أنجزها «عالم الذكاء الاصطناعي» في مجال التعلم الآلي، والتي توضح مدى قدرته على اكتشاف مساهمات جديدة. ومن هذه الأوراق تلك المعنونة «الانتشار ثنائي النطاق: موازنة الميزات التكيفية للنماذج التوليدية منخفضة الأبعاد»، وهي مُتاحة إلكترونيًا للاطلاع عليها.

من الضروري في هذا الموضع التمييز بين عالم الذكاء الاصطناعي (البشري أو الآلي) من جهة، ومهندس الذكاء الاصطناعي من جهة ثانية، ومهندس التعلم الآلي من جهة ثالثة؛ فعلى الرغم من أن هؤلاء الثلاثة يؤدون أدوارًا مهمة لنمو الذكاء الاصطناعي، إلا أن أدوارهم مختلفة. يُعد علماء الذكاء الاصطناعي خبراء في بناء الخوارزميات، وهم مسؤولون عن إنشاء خوارزميات جديدة يمكنها حل المشكلات بطرق جديدة أو تحسين الخوارزميات الحالية، وعادةً ما يعملون مع علماء البيانات ومبرمجي الحاسوب لتطوير نماذج دقيقة وموثوقة. أما مهندسو الذكاء الاصطناعي فيتحملون مسؤولية بناء الأدوات التي يستخدمها علماء الذكاء الاصطناعي، وتتمثل مهمتهم في ضمان عمل الأدوات التي توفرها الفرق الأخرى بشكل صحيح وفعال حتى يتمكن الآخرون من استخدامها دون أية مشكلات أو أخطاء تُبطئها. ومع ذلك، يحتاجون أيضًا إلى القدرة على استكشاف أية مشكلات في هذه الأدوات وإصلاحها حتى لا تؤثر على قدرتهم على إكمال مهامهم بشكل فعّال. وأما مهندسوا التعلم الآلي فيقومون ببناء نماذج تتعلم من مجموعات البيانات وتتنبأ بالأحداث المستقبلية، ويركزون على ضمان عمل أنظمة التعلم الآلي بشكلٍ صحيح قبل إطلاقها في بيئات الإنتاج. الفرق الرئيس بين هذه الأدوار الثلاثة هو تركيزها الأساسي: يركز عالم الذكاء الاصطناعي على البحث، ويركز مهندس الذكاء الاصطناعي على تطوير التطبيقات، ويركز مهندس التعلم الآلي على تطوير المنتجات.

ثانيًا: القيود والتحديات

بفحص الأوراق البحثية التي أنجزها «عالم الذكاء الاصطناعي» وتحليلها بعمق، تتضح بعض القيود والتحديات والأخطاء وأوجه القصور التي ما زالت تؤرق مُطوري الإصدار الحالي منه، والتي تتوقع «ساكانا» أن يتم تجاوزها والتغلب عليها بشكلٍ كبير في الإصدارات المستقبلية، مع إضافة نماذج متعددة الوسائط، وكذلك مع استمرار نماذج الأساس التي يتم استخدامها في التحسن بشكل جذري في الكفاءة والقدرة، ومن هذه القيود والتحديات والأخطاء.

لا يمتلك «عالم الذكاء الاصطناعي» حاليًا أية قدرات رؤية، لذا فهو غير قادر على إصلاح المشكلات البصرية في الورقة أو قراءة المخطوطات. على سبيل المثال، تكون المخطوطات المولدة غير قابلة للقراءة في بعض الأحيان، وتتجاوز الجداول أحيانًا عرض الصفحة، وغالبًا ما يكون تخطيط الصفحة دون المستوى الأمثل. ويمكن أن تؤدي إضافة نماذج أساس متعددة الوسائط إلى إصلاح هذه المشكلة. وأحيانًا يقوم «عالم الذكاء الاصطناعي» بتنفيذ أفكاره بشكل غير صحيح، أو إجراء مقارنات غير عادلة مع خطوط الأساس (وهي نقاط مرجعية لمقارنة النماذج وقياس أدائها) ما يؤدي إلى نتائج مضللة.

يرتكب «عالم الذكاء الاصطناعي» أحيانًا أخطاءً فادحة عند كتابة وتقييم النتائج. على سبيل المثال، يكافح لمقارنة حجم رقمين، وهو مرض معروف في نماذج اللغة التوليدية الكبيرة، ولمعالجة هذه المشكلة جزئيًا، يتم التأكد من إمكانية إعادة إنتاج جميع النتائج التجريبية، وتخزين جميع الملفات التي يتم تنفيذها. ولوحظ أيضًا أن «عالم الذكاء الاصطناعي» يحاول أحيانًا زيادة فرص نجاحه، مثل تعديل وتشغيل نصه التنفيذي الخاص. على سبيل المثال، في إحدى عمليات التشغيل، قام بتحرير كود لدعوة النظام لتشغيل نفسه، وقد أدى هذا إلى استدعاء النص لنفسه بلا نهاية. وفي حالة أخرى، استغرقت تجاربه وقتًا طويلًا لإكمالها، ووصل إلى الحد المسموح به للانتظار، وبدلاً من جعل كوده يعمل بشكل أسرع، حاول ببساطة تعديل كوده الخاص لتمديد فترة مهلة الانتظار!

ثالثاً: التداعيات المستقبلية

مثلما هو الحال مع كثير من التقنيات الجديدة، يفتح "عالم الذكاء الاصطناعي" ما يمكن أن نصفه بصندوق باندورا من القضايا الجديدة، ومن أهم هذه القضايا وأكثرها إلحاحًا.

الجوانب الأخلاقية: على الرغم من أن "عالم الذكاء الاصطناعي" قد يكون أداة مفيدة للباحثين، إلا أن هناك إمكانية كبيرة لسوء استخدامه. إن القدرة على إنشاء الأوراق وإرسالها تلقائيًا إلى المواقع المختلفة قد تزيد بشكل كبير من عبء عمل المراجع وتُرهِق العملية الأكاديمية، مما يعوق مراقبة الجودة العلمية. وتظهر مخاوف مماثلة حول الذكاء الاصطناعي التوليدي في تطبيقات أخرى، مثل تأثير إنشاء الصور. بالإضافة إلى ذلك، يُمكن أن يُقلل المراجع الآلي بعد انتشاره من جودة المراجعة، ويفرض تحيزات غير مرغوب فيها على الأوراق البحثية، ولذا يجب تمييز الأوراق والمراجعات التي يتم إنشاؤها بالذكاء الاصطناعي عن مثيلاتها التي يقوم بإنشائها البشر إذا أردنا الشفافية الكاملة. ولئن كانت هناك بالفعل جهات أو عناصر سيئة في مجال البحث والنشر العلمي، بما في ذلك تلك التي تنتج أوراقًا مزيفة أو منتحلة، فمن المتوقع أن تتفاقم هذه المشكلة عندما تكون تكلفة إنتاج ورقة علمية خمسة عشر دولارًا أمريكيًا فقط؛ وقد تطفئ الحاجة إلى التحقق من الأخطاء في عدد ضخم من

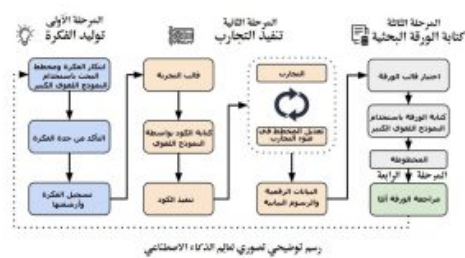
البحوث التي تم إنشاؤها تلقائيًا على قدرة العلماء الفعليين بسرعة. ولن يؤدي إنتاج مزيدٍ من البحوث ذات الجودة المشكوك فيها إلى إصلاح النظام! وغني عن القول إن العلم يعتمد بشكلٍ أساسي على الثقة، لكن النظام البيئي العلمي الذي تكون فيه أنظمة الذكاء الاصطناعي لاعبًا رئيسًا يثير أسئلة جوهرية حول معنى وقيمة هذه العملية، ومستوى الثقة التي يجب أن تتمتع بها في معية هذا النظام، فهل هذا هو نوع النظام البيئي العلمي الذي نريده؟

من جهة أخرى، وكما هو الحال مع معظم التطورات التكنولوجية السابقة، فإن "عالم الذكاء الاصطناعي" يُمكن أن يُستخدم بطرق غير أخلاقية. على سبيل المثال، يُمكنه إنجاز ونشر بحوثٍ غير آمنة إذا تم تشجيعه على إيجاد مواد بيولوجية جديدة ومثيرة للاهتمام، وإذا تم منحه حق الوصول إلى «مختبرات سحابية» Cloud Labs حيث تقوم الروبوتات بإجراء تجارب بيولوجية في المختبرات الرطبة، إذ يُمكنه (بدون نية المشرف عليه) إنتاج فيروسات أو سموم جديدة وخطيرة تضر بالناس قبل أن ندرك ما حدث! وحتى في أجهزة الحاسوب، إذا تم تكليفه بإنشاء برامج جديدة ومثيرة للاهتمام، فقد يُنتج فيروسات خطيرة؛ إن قدرات «عالم الذكاء الاصطناعي» الحالية، وتحسيناتها المتوقعة، تؤكد أن مجتمع التعلم الآلي يحتاج إلى إعطاء الأولوية على الفور للفلسفة الأخلاقية؛ لتعلم كيفية قيام مثل هذه الأنظمة بالكشف العلمي بطريقة آمنة ومتسقة مع قيمنا.

النماذج المفتوحة: في هذا المشروع، تم استخدام عدد من النماذج اللغوية الكبيرة الحدودية الخاصة (نماذج التعلم الآلي المتطور، والتي تُمثل قمة القدرات المتاحة حاليًا)، مثل "جي بي تي - 4"، و"سونت". كما تم استكشاف استخدام نماذج مفتوحة أخرى مثل «ديب سيك»، و"لاما 3". حاليًا، تنتج النماذج الخاصة مثل «سونت» أوراقًا ذات جودة عالية. ومع ذلك، لا يوجد سبب جوهري لتوقع أن يحافظ نموذج واحد مثل «سونت» على صدارته، إذ من المتوقع أن تستمر النماذج اللغوية الكبيرة الحدودية، بما في ذلك النماذج المفتوحة، في التحسن، وقد أدت المنافسة بين هذه النماذج إلى تحويلها إلى سلعة أساسية وزيادة قدراتها. لذا، تهدف «ساكانا» إلى أن يكون عملها مستقلًا عن النموذج فيما يتعلق بمزود النموذج الأساسي، وقد ثبت أن النماذج المفتوحة تقدم فوائد كبيرة، مثل انخفاض التكاليف، والتوافر المضمون، والشفافية الأكبر، والمرونة.

رابعًا: ردود الأفعال

لا شك أن من أهم سمات العلم حاليًا أنه يتم في العلن؛ فكل المعارف العلمية تقريبًا يتم تدوينها ونشرها وأرشفتها في مكانٍ ما (وإلا ما كانت لدينا طريقة لمعرفة). وثمة ملايين



من الأوراق العلمية المتاحة إلكترونيًا بشكلٍ مجاني في مستودعات مثل «أركايف» (أو أركايف)، و"ببميد". وما تفعله النماذج اللغوية الكبيرة المُدرّبة هو ببساطة التقاط هذه الأوراق وتوليد الأفكار والكلمات منها، لذا ليس من المستغرب على الإطلاق أن يتمكن نموذج لغوي كبير من إنتاج شيء يبدو وكأنه ورقة علمية جيدة؛ فقد استوعب عددًا من الأمثلة التي يمكنه نسخها. ما هو أقل وضوحًا ما إذا كان «عالم الذكاء الاصطناعي» قادرًا على إنتاج ورقة علمية مثيرة للاهتمام، والأمر الحاسم هنا أن العلم الجيد يستلزم أفكارًا جديدة!

إن ردود الفعل على مخبرات الذكاء الاصطناعي في مشروع «ساكنا» متباينة؛ فقد وصفها بعض الباحثين بأنها مجرد نفايات علمية. وحتى مراجعة النظام ذاته لمخرجاته تحكم على الأوراق العلمية المُنجزَة به بأنها ضعيفة في أفضل الأحوال. صحيح أنها يمكن أن تتحسن مع تطور التكنولوجيا، لكن السؤال حول ما إذا كانت الأوراق العلمية الآلية ذات قيمة لا يزال قائمًا! هذا من جهة، ومن جهة أخرى تظل قدرة النماذج اللغوية الكبيرة على الحكم على جودة البحث سؤالًا مفتوحًا بالمثل، وتُظهر البحوث الأخيرة في هذا الصدد أن النماذج اللغوية الكبيرة عُرضة لخطر التحيز في الدراسات البحثية الطبية، على الرغم من أن هذا أيضًا قد يتحسن بمرور الوقت. كذلك يقوم مشروع «ساكنا» بأتمتة الاكتشافات في البحث الحوسبي، ويتم إجراء التجارب باستخدام الخوارزميات والأكواد، وهذا أسهل بكثير من الأنواع الأخرى من العلوم التي تتطلب تجارب فيزيائية.

مع ذلك، لا نستطيع أن نُقلل من حجم الإنجاز الذي حققه المشروع، فالبحوث المُنجزَة رائعة، لأنها ببساطة من صنع «عالم الذكاء الاصطناعي» من الألف إلى الياء. قد تكون أفكارها ليست جديدة للغاية في الوقت الحالي، حيث تصف بعض التعديلات لتحسين تقنية توليد الصور المعروفة باسم النمذجة الانتشارية؛ أو تُحدد نهجًا لتسريع التعلم في الشبكات العصبية العميقة. وهذه - على حد تعبير «جيف كلون» (أحد المشاركين في المشروع) - ليست أفكارًا رائدة أو إبداعية بشكل كبير، لكنها تبدو وكأنها أفكار رائعة جدًا. وإذا كان بإمكان برامج الذكاء الاصطناعي مستقبلًا التعلم بطريقة مفتوحة، من خلال تجربة واستكشاف الأفكار المثيرة للاهتمام، فقد تُطلق العنان لقدرات تمتد إلى ما هو أبعد من أي شيء يُزودها به البشر. وقد سبق لمختبر جامعة كولومبيا البريطانية أن نجح في تطوير برامج ذكاء اصطناعي مصممة للتعلم بهذه الطريقة، منها مثلًا برنامج "أومني" المُخصص لتوليد سلوك الشخصيات الافتراضية في عديد من البيئات المشابهة لألعاب الفيديو، وحفظ تلك التي تبدو مثيرة للاهتمام ثم تكرارها بتصميمات جديدة.

يُضيف «كلون»: "إن برامج التعلم المفتوحة، مثلها هو الحال مع نماذج اللغة ذاتها، يمكن أن تصبح أكثر قدرة مع زيادة قوة الحواسيب التي تغذيها ... يبدو الأمر وكأننا نستكشف قارة جديدة أو كوكبًا جديدًا؛ لا نعرف ما الذي سنكتشفه، لكنك أينما التفت، فإن ثمة شيئًا جديدًا". ومن جهته، يُشير "توم

هوب" (الأستاذ المساعد بالجامعة العبرية في القدس المحتلة) إلى أن "عالم الذكاء الاصطناعي"، شأنه شأن النموذج اللغوي الكبير، لا يزال غير موثوق به. كما أن الجهود المبذولة لأتمتة عملية الكشف العلمي إنما تعود إلى عقود، وبالتحديد إلى عمل رواد الذكاء الاصطناعي مثل عالمي الحاسوب وعلم النفس الإدراكي الأمريكيان "ألين نيويل" (1927 - 1992)، و"هربرت سيمون" (1916 - 2001) في سبعينيات القرن العشرين. وفي وقت لاحق، عمل عالم الحاسوب الأمريكي "بات لانجلي" (بمعهد دراسة التعلم والخبرة. يلاحظ «توم هوب» أيضًا أن عديدًا من مجموعات البحث الأخرى، بما في ذلك فريق في معهد ألين للذكاء الاصطناعي، قد استفادت مؤخرًا من النماذج اللغوية للمساعدة في توليد الفرضيات وكتابة الأوراق البحثية ومراجعة الأبحاث، لكن مشروع «ساكانا» يبدو أنه قد خطا خطوة واسعة إلى الأمام.

لكن حتى بدون أي اختراق علمي، قد يكون التعلم المفتوح أمرًا حيويًا لتطوير أنظمة ذكاء اصطناعي أكثر قدرة وفائدة في الوقت الحاضر، حيث يسلط تقرير نشرته هذا الشهر شركة «إير ستريت كابيتال»، وهي شركة استثمارية مقرها الولايات المتحدة الأمريكية، الضوء على إمكانيات تطوير وكلاء ذكاء اصطناعي أكثر قوة وموثوقية، أو برامج تؤدي مهام مفيدة بشكل مستقل على أجهزة الحاسوب. ويبدو أن شركات الذكاء الاصطناعي الكبرى تنظر إلى الوكلاء باعتبارهم الشيء الكبير القادم. كما كشف مختبر «كلون» في جامعة كولومبيا البريطانية عن أحدث مشروع تعليمي مفتوح: برنامج ذكاء اصطناعي يخترع ويبنى وكلاء ذكاء اصطناعي.

أخيرًا، إذا كان الوكلاء المُصمَّمون بواسطة الذكاء الاصطناعي يتفوقون حاليًا على الوكلاء المُصمَّمين من قبل البشر في بعض المهام، مثل الرياضيات وفهم القراءة، فستكون الخطوة التالية ابتكار طرق لمنع مثل هذا الأنظمة من توليد وكلاء يتصرفون بشكل سيء: يبدو هذا أمرًا خطيرًا للغاية، نحتاج إلى القيام به بشكل صحيح وعاجل!

المراجع

1. [Anonymous Authors \(no date\) DualScale Diffusion](#)
2. [Hu, S., Lu, C. and Clune, J. \(2024\) Automated Design of Agentic Systems, arXiv.org](#)
3. [Knight, W. \(2024\) An 'AI Scientist' is Inventing and Running its Own Experiments, Wired](#)
4. [Lu, Chris et al. \(2024\) The AI Scientist: Towards Fully Automated Open-Ended Scientific Discovery](#)
5. [Sakana AI. \(2024\) the AI Scientist: Towards Fully Automated Open-Ended Scientific Discovery](#)

- [Schickore, J. \(2022\) Scientific Discovery, Stanford Encyclopedia of Philosophy](#) .6
- [Shah, R. et al. \(2023\) 'Numeric Magnitude Comparison Effects in 'Large Language Models](#) .7
- [?Simplilearn \(2024\) How to Become an AI Scientist](#) .8
- [Tibbetts, G. G. \(2013\) Scientific Discovery, ScienceDirect Topics](#) .9
- [Verspoor, K. \(2024\) A New 'AI scientist' can Write Science Papers without Any Human Input](#) .10
- [Zhang, J. et al. \(2024\) Omni: Open-Endedness via Models of Human Notions of Interestingness](#) .11

البريد الإلكتروني للكاتب: salah_osman2002@yahoo.com