

منصات خوارزمية جديدة للتنقيب في البيانات الطبية

استخدامها في علاج كوفيد – 19، أو العلاقات الجينية التي يمكن أن تسرع من الوصول إلى لقاح أو نظام علاجي فعال. وهو ما يساهم به محرك بحث أطلقه مختبر لورنس بيركلي الوطني منذ أيام باعتماد أحدث تقنيات الذكاء الاصطناعي، متيحاً للعلماء أدوات تنقيب عن الأدوية وتأثيراتها لتسريع أبحاثهم حول كوفيد.

الباحثون يأملون أن يساعد الذكاء الاصطناعي ومعالجة اللغات في اكتشاف المعرفة العلمية الكامنة في نصوص المقالات

ويشير الباحثون إلى أن أهم عناصر منصتهم الجديدة كوفيد سكولار (COVIDSCHOLAR) هو البيانات، حيث قاموا ببناء برمجيات تجمع الأوراق البحثية من مجموعة كبيرة من المصادر لإتحاتها على المنصة في غضون 15 دقيقة من وقت نشرها على الإنترنت. كما تولي المنصة مهمة إصلاح أخطاء التنسيق، ثم يبدأ عمل خوارزميات التعلم الآلي التي تضيف الوسوم المناسبة لكل ورقة بحثية لتساعد في تصنيفها إلى فئات بحسب موضوع البحث وأهميته.

وحتى الآن، تتضمن المنصة أكثر من 60 ألف ورقة بحثية حول كوفيد – 19، وتستمر في التوسع بشكل يومي. ويطور الباحثون خوارزميات الذكاء الاصطناعي حتى تسمح للعلماء بتنظيم نتائج البحث في نماذج كمية لدراسة مواضيع محددة مثل التفاعلات بين البروتينات. كما عمل الباحثون على إضافة الدراسات العلمية السابقة إلى قاعدة البيانات من أجل اكتشاف أي معلومات هامة قد تساعد في مواجهة الوباء.

ويؤكد القائمون على المنصة أن الذكاء الاصطناعي لن يحل محل العلماء، لكنه يساعدهم في تسريع جهودهم لإيجاد الترياق المناسب لكوفيد – 19.

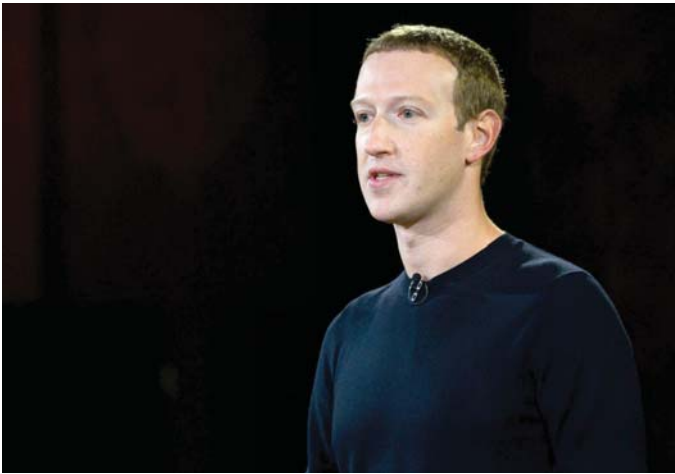
فيسبوك تقرب بين الواقعين الافتراضي والحقيقي

(VR)– مفيدة في الحياة اليومية مثل الهواثف الذكية“.

حاول المشروع استخدام الأسلوب البشري في التذكر؛ على غرار التساؤل وتذكر الأماكن التي قد أضاع أحدنا فيها مفاتيح سيارته، أو تذكر موقف أو جملة معينة قالها البعض في موقف ما، أو التركيز على كلام صديق في بيئة صاخبة.

تقول كريستين جرومان، كبيرة علماء الأبحاث في فيسبوك، ”تقليدياً يتعلم النظام الجديد من خلال محاولة القيام بأشياء لاحتفظها في العالم، أو دفعه إلى تعلم كيفية القيام بامور معينة“.

وعلى الرغم من أن هذه المهام لا يمكن القيام بها بواسطة أي نظام ذكاء اصطناعي في الوقت الحالي، إلا أنها قد تكون جزءاً كبيراً من خطط فيسبوك للجمع بين الحقيقة والواقع الافتراضي. وفي يوليو الماضي كشف الرئيس التنفيذي مارك زوكربيرغ النقاب عن خطط فيسبوك لجعل نظارات الواقع المعزز هي الطفرة الثانية التي سيعتمد عليها مستقبلاً لتصفح الإنترنت.



زوكربيرغ يعد بطفرة في تصفح الإنترنت

لندن – نتائج الدراسات الطبية تشكل ثروة هائلة من البيانات، إلا أنها موزعة عبر مجموعات مختلفة، مما يترك العديد من الأسئلة الطبية دون إجابة. مثال على ذلك، في حال أظهرت مجموعة بيانات علاقة بين السمعة وأمراض القلب، وأظهرت دراسة أخرى وجود علاقة بين انخفاض فيتامين (د) والسمعة، فما هي العلاقة بين انخفاض فيتامين (د) وأمراض القلب؛ سيتطلب إيجاد هذه العلاقة تجربة سريرية أخرى.

كيف يمكننا استغلال هذه المعلومات الجزئية بشكل أفضل؟ تستطيع أجهزة الكمبيوتر اكتشاف الأنماط المختلفة بسهولة، ولكن يُعتبر الارتباط بين هذه الأنماط في الطب ارتباطاً عادياً، وليس سببياً. وفي السنوات الماضية، بدأ العلماء ببرمجة بعض الخوارزميات القادرة على تحديد العلاقات السببية داخل مجموعة بيانات واحدة. وهنا تكمن المُشكلة، حيث يُعتبر البحث في كل مجموعة بيانات بشكل فردي أمراً صعباً للغاية، فهو يشبه البحث في عدد كبير من الأقفال واحداً تلو الآخر. أي أننا نحتاج إلى إيجاد طريقة تمكننا من البحث في جميع المجموعات معاً.

وهنا توصل الباحثان أنيش دهير وسياران لي من شركة بايبلون هيلث – وهي شركة متخصصة في توفير الرعاية الصحية الرقمية في المملكة المتحدة – إلى تقنية لإيجاد العلاقات السببية عبر مجموعات البيانات المختلفة. مما يفتح المجال لاستخدام نوع جديد من قواعد البيانات الطبية الكبيرة في البحث عن الأسباب والنتائج، التي لم نستطع استخدامها من قبل، مما يؤدي إلى زيادة احتمال اكتشافنا لروابط سببية جديدة. ويأمل الباحثون أن يساعد الذكاء الاصطناعي ومعالجة اللغات الطبيعية في اكتشاف المعرفة العلمية الكامنة في نصوص المقالات من خلال النقاط الارتباطات التي تغيب عن العقل البشري والتي يتطلب إيجادها جمع المعلومات من مئات آلاف الأبحاث العلمية.

على سبيل المثال، يمكن للذكاء الاصطناعي أن يساهم في اكتشاف الأدوية الموجودة التي يمكن

الشبكات العصبية العميقة والأدمغة.. اختلافات منهجية وأوجه تشابه



الذكاء الاصطناعي... يسمع يرى ويتكلم ولكن هل يتخذ القرار؟

الحسابات المعقدة تافهة بالنسبة لها، إلا أن بعض المهام التي يسهل حلها على البشر نسبياً، يصعب إكمالها على هذه الشبكات.

درس الفريق 13 تأثيراً إيجابياً مختلفاً وكشف عن اختلافات نوعية لم تكن معروفة سابقاً بين الشبكات العميقة والدماغ البشري. مثال على ذلك هو ما يعرف بتأثير ”تأثير“، وهي ظاهرة يجد فيها البشر أنه من الأسهل التعرف على تغيرات السمات المحلية في صورة عمودية، ولكن هذا يصبح صعباً عندما تنقلب الصورة رأساً على عقب.

أظهرت الشبكات العميقة المدربة على التعرف على الوجوه المستقيمة تأثير تأتشر عند مقارنتها بالشبكات المدربة على التعرف على الأشياء. تم اختبار خاصية بصرية أخرى للدماغ البشري، تسمى ارتباط المرأة، على هذه الشبكات. بالنسبة إلى البشر، تبدو الانعكاسات المرأة على طول المحور الراسي أكثر تشابهاً من تلك الموجودة على طول المحور الأفقي. وجد الباحثون أن الشبكات العميقة تظهر أيضاً ارتباطاً أقوى في المرأة للصور الرأسية مقارنة بالصور المنعكسة أفقياً.

ظاهرة أخرى خاصة بالدماغ البشري، وهي التركيز على الشكل العام أولاً. يعرف هذا بتأثير الميزة العالمية، على سبيل المثال، في صورة شجرة، سيرى دماغنا أولاً الشجرة ككل قبل ملاحظة تفاصيل الأوراق فيها. وبالمثل، عند تقديم صورة للوجه، ينظر البشر أولاً إلى الوجه ككل، ثم ينتقلون إلى التركيز على التفاصيل الدقيقة مثل العينين والأنف والفم وإلى ذلك، كما يوضح جورجين جاكوب، مشارك في الدراسة وطالب دكتوراه في المعهد، حيث يقول ”من المثير للدهشة أن الشبكات العصبية أظهرت ميزة محلية“. هذا يعني أنه على عكس الدماغ، تركز الشبكات على التفاصيل الدقيقة للصورة أولاً. لذلك، على الرغم من أن هذه الشبكات العصبية والدماغ البشري ينفذان نفس مهام التعرف على الأشياء، فإن الخطوات التي يتبعها الاثنان مختلفة تماماً.

يقول أرون، وهو كبير مؤلفي الدراسة ”أظهرت الكثير من الدراسات أوجه التشابه بين الشبكات العميقة والأدمغة، لكن لم ينظر أحد حقاً في الاختلافات المنهجية“.

يمكن أن يدفعنا تحديد هذه الاختلافات إلى جعل هذه الشبكات أقرب إلى الدماغ، ويمكن لمثل هذه التحليلات أن تساعد الباحثين على بناء شبكات عصبية أكثر قوة لا تؤدي فقط بشكل أفضل ولكنها أيضاً محصنة ضد ”الهجمات العدائية“ التي تهدف إلى إخراجها عن مسارها.

ويأمل الباحثون أن يمدد هذا العمل الطريق لتقنية ذكاء اصطناعي أكثر موثوقية تتصرف مثل البشر وتقليل الأخطاء التي لا يمكن التنبؤ بها مستقبلاً.

في حدود اثنتين. تضمن كل مقطع أيضاً ضوضاء في الخلفية لجعل المهمة أكثر واقعية (وأكثر صعوبة). بعد عدة آلاف من الأمثلة، تعلم النموذج أداء المهمة بنفس دقة المستمع البشري. والفكرة هي أنه بمرور الوقت يصبح النموذج أفضل وأفضل في المهمة، والأمل هو أن يتعلم شيئاً عاماً.

ظاهرة ”تأتشر“

كان النموذج يميل أيضاً إلى ارتكاب الأخطاء في نفس المقاطع التي ارتكبتها البشر معظم الأخطاء. وفي دراسة حديثة أخرى، قارن SP (ARUN)، الأستاذ المساعد في CNS، وفريقه الخصائص النوعية المختلفة لهذه الشبكات العميقة مع تلك الموجودة في الدماغ البشري.

وفي الدراسة الحالية المنشورة في نشرة ”ناتشرل كومونيكيشن“ للاتصالات الطبيعية، حاول أرون وفريقه فهم المهام المرئية التي يمكن أن تؤديها هذه الشبكات بشكل طبيعي بحكم هندستها المعمارية، والتي تتطلب المزيد من التدريب.

وعلى الرغم من أن الشبكات العميقة نموذج جيد لفهم كيفية تصور الدماغ البشري للأشياء، إلا أنها تعمل بشكل مختلف عن دماغ البشر. وفي حين أن



كانت هذه فرصة مثيرة لعلم الأعصاب، حيث تمكن العلماء من إنشاء أنظمة يمكنها فعل بعض الأشياء التي يمكن للناس القيام بها، ومن ثم استجواب النماذج ومقارنتها بالدماغ. وفي هذا الصدد قام باحثو معهد ماساتشوستس للتكنولوجيا بتدريب شبكات عصبية على

أداء مهمتين سمعيتين، إحداهما تتضمن الكلام والأخرى تتضمن الموسيقى. بالنسبة إلى مهمة الكلام، قدم الباحثون للنموذج الآلاف من التسجيلات لمدة ثابنتين لشخص يتحدث. كانت المهمة تحديد الكلمة في منتصف المقطع. بالنسبة إلى مهمة الموسيقى، طلب من النموذج تحديد نوع مقطع موسيقي مدته

من الوجهة التي يستخدمها إنسان آخر للتعرف عليه. إذا لم يفعل الذكاء الاصطناعي هذا، فقد يكون لدينا وهم بأن النظام يعمل تماماً مثل البشر، ولكن بعد ذلك نجد أنه يخطئ في بعض الظروف الجديدة أو غير المختبرة“.

واستخدم الباحثون سلسلة من الوجوه ثلاثية الأبعاد القابلة للتعديل، وطلبوا من البشر تقييم تشابه هذه الوجوه التي تم إنشاؤها عشوائياً إلى أربع هويات مألوفة. ثم استخدموا هذه المعلومات لاختبار ما إذا كانت الشبكات العصبية العميقة قد صنعت نفس التصنيفات لنفس الأسباب – اختبار ليس فقط ما إذا كان البشر والذكاء الاصطناعي قد اتخذوا نفس القرارات، ولكن أيضاً ما

إذا كانت تستند إلى نفس المعلومات. الأهم من ذلك، من خلال نهجهم هذا، يمكن للباحثين تصور هذه النتائج على أنها وجوه ثلاثية الأبعاد تقود سلوك البشر والشبكات. على سبيل المثال، الشبكة التي صنفت 2000 هوية بشكل صحيح كانت مدفوعة بوجه كاركاتيري بشكل كبير، مما يدل على أنها حدثت الوجوه التي تعالج معلومات وجه مختلفة تماماً عن البشر.

نمذجة الدماغ البشري

عندما تم تطوير الشبكات العصبية العميقة لأول مرة في الثمانينات، كان علماء الأعصاب يأملون في إمكانية استخدام مثل هذه الأنظمة لنمذجة الدماغ البشري. ومع ذلك، لم تكن أجهزة الكمبيوتر من تلك الحقبة قوية بما يكفي لبناء نماذج كبيرة بما يكفي لأداء مهام العالم الحقيقي مثل التعرف على الأشياء

أو التعرف على الكلام. على مدى السنوات الخمس الماضية، اتاحت التطورات في قوة الحوسبة وتكنولوجيا الشبكات العصبية استخدام الشبكات العصبية لأداء مهام صعبة في العالم الحقيقي، وأصبحت هي النهج القياسي في العديد من التطبيقات الهندسية. بالتوازي مع ذلك، أعاد بعض علماء الأعصاب النظر في إمكانية استخدام هذه الأنظمة لنمذجة الدماغ البشري.

كانت هذه فرصة مثيرة لعلم الأعصاب، حيث تمكن العلماء من إنشاء أنظمة يمكنها فعل بعض الأشياء التي يمكن للناس القيام بها، ومن ثم استجواب النماذج ومقارنتها بالدماغ. وفي هذا الصدد قام باحثو معهد ماساتشوستس للتكنولوجيا بتدريب شبكات عصبية على أداء مهمتين سمعيتين، إحداهما تتضمن الكلام والأخرى تتضمن الموسيقى. بالنسبة إلى مهمة الكلام، قدم الباحثون للنموذج الآلاف من التسجيلات لمدة ثابنتين لشخص يتحدث. كانت المهمة تحديد الكلمة في منتصف المقطع. بالنسبة إلى مهمة الموسيقى، طلب من النموذج تحديد نوع مقطع موسيقي مدته

إنشاء ذكاء اصطناعي مشابه لذكاء الإنسان هو أكثر من مجرد محاكاة لسلوك البشر: يجب أن تكون التكنولوجيا أيضاً قادرة على معالجة المعلومات، أو ”التفكير“، بطريقة تماثل طريقة البشر في معالجة المعلومات إذا كان الهدف الاعتماد عليها دون تدخل من العنصر البشري.

لندن – في بحث جديد، نُشر في مجلة ”باتيرنز“ بإشراف من القائمين على كلية علم النفس وعلم الأعصاب بجامعة غلاسكو، يستخدم النمذجة ثلاثية الأبعاد لتحليل الطريقة التي تتبناها الشبكات العصبية العميقة – وهي جزء من العائلة الأوسع للتعلم الآلي – في معالجة المعلومات، لتصور كيفية معالجة المعلومات الخاصة بهم يطابق البشر.

ويأمل القائمون على الدراسة أن يمدد هذا العمل الجديد الطريق لإنشاء تقنية ذكاء اصطناعي يمكن الاعتماد عليها، ويمكن بواسطتها معالجة المعلومات بطريقة تماثل طرق البشر، وبالتالي أن ترتكب أخطاء يمكننا فهمها والتنبؤ بها.

تناقضات وأخطاء

يتمثل أحد التحديات التي لا تزال تواجه تطوير الذكاء الاصطناعي في كيفية فهم عملية التفكير الآلي بشكل أفضل، وما إذا كانت تتطابق مع كيفية معالجة البشر للمعلومات، من أجل ضمان الدقة. وغالباً ما يتم تقديم الشبكات العصبية العميقة على أنها أفضل نموذج حالي لسلوك صنع القرار البشري، حيث تحقق أوا حتى تتجاوز الأداء البشري في بعض المهام. ومع ذلك، حتى مهام التمييز المرعي البسيطة المخادعة يمكن أن تكشف عن تناقضات وأخطاء واضحة من نماذج الذكاء الاصطناعي، عند مقارنتها بالبشر.

حالياً، تُستخدم تقنية الشبكات العصبية العميقة في تطبيقات مثل التعرف على الوجوه، وعلى الرغم من أنها ناجحة جداً في هذه المجالات، لا يزال العلماء عاجزون تماماً عن فهم كيفية معالجة هذه الشبكات للمعلومات، وبالتالي عاجزون أيضاً عن تقديم تفسير لحدوث أخطاء.

والشبكات العصبية العميقة هي أنظمة تعلم آلي مستوحاة من شبكة خلايا الدماغ أو الخلايا العصبية في الدماغ البشري، والتي يمكن تدريبها لأداء مهام محددة. ولعبت هذه الشبكات دوراً محورياً في مساعدة العلماء على فهم كيفية إدراك ادماغنا للأشياء التي نراها.

على الرغم من أن الشبكات العميقة قد تطورت بشكل كبير خلال العقد الماضي، إلا أنها لا تزال بعيدة كل البعد عن أداء الدماغ البشري في إدراك الإشارات البصرية.

في الدراسة الجديدة، عالج فريق البحث هذه المشكلة من خلال نمذجة المحفز البصري الذي أعطته الشبكة العصبية العميقة، ومن ثم تحويله بطرق متعددة حتى يتمكنوا من إظهار تشابه في التعرف من خلال معالجة المعلومات المماثلة بين البشر ونموذج الذكاء الاصطناعي.

على الرغم من أن الشبكات العميقة نموذج جيد لفهم تصور الدماغ للأشياء إلا أنها تعمل بشكل مختلف عن دماغ البشر

قال البروفيسور فيليب شينز، كبير مؤلفي الدراسة ورئيس معهد علوم الأعصاب والتكنولوجيا بجامعة غلاسكو ”عند بناء نماذج ذكاء اصطناعي تتصرف مثل البشر، على سبيل المثال التعرف على وجه الشخص متى رأى أنه إنسان، يجب أن نتأكد من أن نموذج الذكاء الاصطناعي يستخدم نفس المعلومات