



مقدمة قصيرة جداً

ما جريت إيه بهدين

الذكاء الاصطناعي

ترجمة إبراهيم سند أحمد

الذكاء الاصطناعي

مقدمة قصيرة جدًا

تأليف
مارجريت إيه بودين

ترجمة
إبراهيم سند أحمد

مراجعة
هاني فتحي سليمان



الناشر مؤسسة هنداوي

المشهرة برقم ١٠٥٨٥٩٧٠ بتاريخ ٢٦ / ١ / ٢٠١٧

يورك هاوس، شيت ستريت، وندسور، SL4 1DD، المملكة المتحدة

تليفون: ٨٣٢٥٢٢ ١٧٥٣ (٠) ٤٤ +

البريد الإلكتروني: hindawi@hindawi.org

الموقع الإلكتروني: https://www.hindawi.org

إنَّ مؤسسة هنداوي غير مسئولة عن آراء المؤلف وأفكاره، وإنما يعبر الكتاب عن آراء مؤلفه.

تصميم الغلاف: ولاء الشاهد

الترقيم الدولي: ٩٧٨ ١ ٥٢٧٣ ٢٩٩٠ ٤

صدر الكتاب الأصلي باللغة الإنجليزية عام ٢٠١٨.

صدرت هذه الترجمة عن مؤسسة هنداوي عام ٢٠٢٢.

جميع حقوق النشر الخاصة بتصميم هذا الكتاب وتصميم الغلاف محفوظة لمؤسسة هنداوي.

جميع حقوق النشر الخاصة بالترجمة العربية لنص هذا الكتاب محفوظة لمؤسسة هنداوي.

جميع حقوق النشر الخاصة بنص العمل الأصلي محفوظة لدار نشر جامعة أكسفورد.

Copyright © Margaret A. Boden 2016, 2018. *Artificial Intelligence* was originally published in English in 2016, 2018. This translation is published by arrangement with Oxford University Press. Hindawi Foundation is solely responsible for this translation from the original work and Oxford University Press shall have no liability for any errors, omissions or inaccuracies or ambiguities in such translation or for any losses caused by reliance thereon.

المحتويات

٩	شكر وتقدير
١١	١- ما الذكاء الاصطناعي؟
٢٧	٢- كَأَن الذكاء العام هو الكأس المقدسة
٥٧	٣- اللغة، الإبداع، العاطفة
٧٥	٤- الشبكات العصبية الاصطناعية
٩٥	٥- الروبوتات والحياة الاصطناعية
١٠٩	٦- لكن هل هو ذكاء حَقًّا؟
١٣١	٧- التفرد
١٥١	قراءات إضافية
١٥٣	المراجع
١٦١	مصادر الصور

إهداء إلى بايرون، وأوسكار، ولوكاس، وألينا

شكر وتقدير

أُتوجّه بالشكر إلى أصدقائي الآتية أسماؤهم على كل مشورة مفيدة قدّموها لي (مع تحملي الكامل للمسئولية عن أي أخطاء بطبيعة الحال): فيل هاسباندر وجيرمي ريفين وأنيل سيث وأرون سلومان وبلاي ويتبي. كما أتقدّم بالشكر إلى لاثا مينون على تفهّمها وتحليلها بالصبر.

الفصل الأول

ما الذكاء الاصطناعي؟

الهدف الأساسي من الذكاء الاصطناعي هو تمكين أجهزة الكمبيوتر من تنفيذ المهام التي يستطيع العقل تنفيذها.

عادةً ما يُطلق على بعض تلك المهام (مثل التفكير) صفة «الذكاء». وبعضها (مثل الرؤية) لا يُطلق عليه ذلك الوصف. ولكن جميعها لا يخلو من مهاراتٍ نفسية تمكّن الإنسان والحيوان من الوصول إلى أهدافهما، ومن تلك المهارات الإدراك الحسي، والربط بين الأفكار، والتنبؤ، والتخطيط، والتحكم الحركي.

لا ينطوي الذكاء على بُعد واحد، ولكنه مساحة غنية بالتنظيم، وتضم قدرات متنوعة لمعالجة المعلومات. ومن ثم، يستخدم الذكاء الاصطناعي العديد من التقنيات المختلفة التي تنفّذ العديد من المهام المختلفة.

أضف إلى ذلك أن الذكاء الاصطناعي موجود في كل مكان.

توجد الاستخدامات العملية للذكاء الاصطناعي في المنازل، والسيارات (والسيارات بدون سائق)، والمكاتب، والبنوك، والمستشفيات، والفضاء ... وشبكة الإنترنت، بما في ذلك إنترنت الأشياء (الذي يربط المستشعرات المادية التي يتزايد استخدامها في الأجهزة والملابس والبيئات). بعض تلك الاستخدامات يكون خارج الكوكب، مثل الروبوتات التي تُرسل إلى القمر والمريخ، أو الأقمار الصناعية التي تدور في الفضاء. أفلام الرسوم المتحركة في هوليوود، وألعاب الفيديو والكمبيوتر، وأنظمة الملاحة عبر الأقمار الصناعية، ومحركات بحث «جوجل»، جميعها تعتمد على تقنيات الذكاء الاصطناعي. ومن ذلك أيضاً الأنظمة التي يستخدمها المستثمرون للتنبؤ بتحركات البورصة والأنظمة التي تستخدمها الحكومات الوطنية للإسهام في توجيه القرارات المتعلقة بشأن الصحة والنقل والمواصلات. ومن ذلك أيضاً التطبيقات على الهواتف المحمولة. أضف إلى ذلك الصور الرمزية في الواقع الافتراضي

والنماذج التجريبية للعاطفة التي تُطوّر في الروبوتات «المرافقة». حتى المعارض الفنية تستخدم الذكاء الاصطناعي على مواقعها الإلكترونية، وحتى في معارض الفنون الحاسوبية. لحسن الحظ أن الطائرات العسكرية التي بدون طيار تجول في ساحات القتال اليوم، ولحسن الحظ أكثر أن كاسحات الألغام الروبوتية أيضاً تفعل ذلك.

ثمة هدفان أساسيان للذكاء الاصطناعي. الهدف الأول «تكنولوجي»؛ استخدام أجهزة الكمبيوتر لإنجاز مهام مفيدة (وتوظّف في بعض الأحيان طرقاً غير التي يستخدمها العقل تماماً). الهدف الثاني «علمي»؛ استخدام مفاهيم الذكاء الاصطناعي ونماذجها للمساعدة في الإجابة عن أسئلة تتعلق بالإنسان وغيره من الكائنات الحية. لا يركّز معظم العاملين في الذكاء الاصطناعي إلا على هدف من هذين الهدفين، ولكن بعضهم يركّز على كليهما.

يوفّر الذكاء الاصطناعي عدداً لا يُحصى من الأدوات التكنولوجية، أضف إلى ذلك تأثيره العميق في علوم الحياة.

وعلى وجه التحديد، مكّن الذكاء الاصطناعي علماء النفس وعلماء الأعصاب من وضع نظريات راسخة عن العقل والدماع. تتضمن تلك النظريات نماذج عن «آلية عمل دماغ الإنسان»، وسؤالاً آخر لا يقل أهمية وهو: «ما الذي يفعله الدماغ؟» ما الأسئلة الحاسوبية (السيكولوجية) التي يجيب عنها؟ وما أنواع معالجة المعلومات التي تمكّنه من فعل ذلك؟ هناك العديد من الأسئلة التي لم يُجب عنها، وقد علّمنا الذكاء الاصطناعي نفسه أن عقولنا أغنى مما تصوّره علماء النفس من قبل.

كذلك استخدم علماء الأحياء الذكاء الاصطناعي في صورة «حياة اصطناعية» تطوّر نماذج حاسوبية لجوانب مختلفة لدى الكائنات الحية. وهذا يساعدهم في شرح الأنواع المختلفة من سلوك الحيوان، ونمو الجسم، والتطوّر الأحيائي، وطبيعة الحياة نفسها.

بالإضافة إلى التأثير في علوم الحياة، أثر الذكاء الاصطناعي في الفلسفة. يؤسّس العديد من الفلاسفة اليوم أفكارهم على مفاهيم الذكاء الاصطناعي. على سبيل المثال، يستخدم الفلاسفة تلك المفاهيم لحل المعضلة الشهيرة بين العقل والجسم، ولغز الإرادة الحرة، والعديد من الألغاز المتعلقة بالوعي. لكن هذه الأفكار الفلسفية جدلية إلى حدّ كبير. ويدور خلاف كبير بشأن ما إذا كان هناك أي نظام للذكاء الاصطناعي يمتلك شكلاً حقيقياً من الذكاء أو الإبداع أو الحياة.

وأخيراً وليس آخراً، تحدّى الذكاء الاصطناعي الطرق التي نفكر بها بشأن الإنسانية ومستقبلها. في الحقيقة، يخشى البعض ولا يدري هل المستقبل لنا حقاً أم لا؛ لأنهم

يستشرفون تفوق الذكاء الاصطناعي على الذكاء البشري في شتى المجالات. وعلى الرغم من ترحيب بعض المفكرين بهذا المستقبل فإن الغالبية تخشاه؛ إذ يتساءلون: ماذا سيبقى لكرامة الإنسان ومسئوليته؟ سنتناول كل تلك المسائل في الفصول الآتية.

الأجهزة الافتراضية

قد يقول أحدهم: «الحديث عن الذكاء الاصطناعي يعني الحديث عن أجهزة الكمبيوتر». الإجابة لها وجهان. ولكن أجهزة الكمبيوتر ليست محور الموضوع. بل إنها أدوات تُستخدم في هذا الشأن. بعبارة أخرى، على الرغم من أن الذكاء الاصطناعي بحاجة إلى أجهزة «مادية» (مثل أجهزة الكمبيوتر)، فالأحرى أن تنصرف أذهاننا إلى ما يُسمّى علماء الكمبيوتر «الأجهزة الافتراضية».

الجهاز الافتراضي ليس جهازاً مصوراً في الواقع الافتراضي، وليس كتقليد محرّك السيارة المُستخدَم في تدريب الميكانيكيين. بل هو «نظام لمعالجة المعلومات» يتصوره المُبرمج في عقله عندما يكتب برنامجاً، ويتصوره الناس في عقولهم عندما يستخدمونه. على سبيل المثال، فُكر المصمّم في تطوير مُعالِج كلمات وجربّه المستخدمون ممّن لهم تعامل مباشر مع الكلمات والفقرات. ولكن البرنامج نفسه ليس من مكوّناته الفكرة أو التجربة. يُعتقد كذلك أن الشبكة العصبية (انظر الفصل الرابع) تعالج المعلومات بـ «التوازي»، على الرغم من أنه عادةً ما تُنفَّذ باستخدام جهاز كمبيوتر (تسلسلي) بهيكله فون نيومان.

هذا لا يعني أن الجهاز الافتراضي مجرد قصة نرويها أو مجرد ضرب من الخيال. الأجهزة الافتراضية حقائق واقعية. بإمكان تلك الأجهزة تنفيذ مهمات، سواء داخل النظام أو في العالم الخارجي (إذا رُبِطت بأجهزة مادية مثل الكاميرا أو أيدي الروبوت). يحاول العاملون في مجال الذكاء الاصطناعي أن يكتشفوا مَكَمَن الخطأ عندما يفعل البرنامج شيئاً غير متوقَّع، ونادراً ما يأخذون في الاعتبار أعطال مكونات الأجهزة. وعادةً ما يهتمون بالأحداث والتفاعلات السببية في الجهاز الافتراضي أو البرنامج.

لغات البرمجة كذلك عبارة عن أجهزة افتراضية (حيث تُترجم تعليماتها إلى تعليمات برمجية للجهاز قبل تشغيلها). تُعرف بعض لغات البرمجة بأنها لغات منخفضة المستوى؛ ومن ثَمّ تلزم الترجمة على عدة مستويات.

لا يقتصر الأمر على لغات البرمجة وحدها. تتكوّن الأجهزة الافتراضية بوجه عام من أنماط أنشطة (معالجة المعلومات) ذات عدة مستويات. إضافة إلى ذلك، ليست الأجهزة الافتراضية وحدها هي التي تعمل على أجهزة الكمبيوتر. سنرى في الفصل السادس أنه يمكن فهم «العقل البشري» وكأنه جهاز ظاهري مُركَّب في الدماغ، أو بالأحرى مجموعة من الأجهزة الافتراضية التي تتبادل التفاعل وتعمل بالتوازي بعضها مع بعض (وقد تطوّرت أو تعلّمت في أوقات مختلفة).

التقدّم في الذكاء الاصطناعي يتطلب التقدّم في تحديد الأجهزة الافتراضية المهمة/المفيدة. كلما زادت قوة أجهزة الكمبيوتر مادياً (من حيث الحجم أو السرعة)، كان ذلك أفضل. وربما تكون القوة ضرورية لأنواع مُعيّنة من الأجهزة الافتراضية المراد تنفيذها. ولكن لا يمكن استخدام أجهزة الكمبيوتر إلا إذا أمكن تشغيل الأجهزة الافتراضية القوية في معالجة المعلومات. (وبالمثل، التقدّم في علم الأعصاب يتطلب فهماً أفضل للأجهزة الافتراضية النفسية التي تُنفّذها الخلايا العصبية المادية؛ انظر الفصل السابع).

نُستخدم أنواع مختلفة من معلومات العالم الخارجي. فكل نظام يعمل بالذكاء الاصطناعي يحتاج إلى أجهزة مُدخلات ومُخرجات، حتى ولو لوحة مفاتيح وشاشة. وفي كثير من الأحيان، توجد مستشعرات تُستخدم لأغراض خاصة (ربما كاميرات أو شعيرات حساسة للضغط) و/أو أجهزة استجابة (ربما معدات تركيب الصوت للموسيقى أو الكلام، أو أذرع الروبوت). يتصل برنامج الذكاء الاصطناعي بأجهزة الكمبيوتر تلك، ويُحدث تغييرات في الواجهات العالمية، وكذلك يعالج المعلومات داخلياً.

عادةً ما تتضمن معالجة الذكاء الاصطناعي أجهزة مدخلات ومخرجات «داخلية»، بحيث تمكّن الأجهزة الافتراضية المتنوعة داخل النظام بأكمله من التفاعل بعضها مع بعض. على سبيل المثال، ربما يكتشف جزء في برنامج الشطرنج تهديداً محتملاً برصد شيء يحدث في جزء آخر، وربما يربطه عندئذٍ بجزء ثالث بحثاً عن حركة حظر.

أنواع الذكاء الاصطناعي الأساسية

تعتمد طريقة معالجة المعلومات على الجهاز الافتراضي المستخدم. وكما سنرى في الفصول اللاحقة، يوجد خمسة أنواع أساسية، وكل نوع يضم العديد من التباينات. النوع الأول هو الذكاء الاصطناعي الكلاسيكي — أو الرمزي — ويُطلق عليه في بعض الأحيان الذكاء الاصطناعي التقليدي الجميل. يوجد نوع آخر، وهو الشبكات العصبية الاصطناعية أو

الترابطية. إضافةً إلى ذلك، يوجد نوع من البرمجة التطورية، وهو الأتمتة الخلوية، والأنظمة الديناميكية.

غالبًا ما يستخدم فردى الباحثين طريقةً واحدة، ولكن هذا لا ينفي استخدام أجهزة افتراضية «مختلطة». على سبيل المثال، ورد في الفصل الرابع نظرية الفعل البشري الذي لا يتوقف عن التبديل بين المعالجة الرمزية والاتصالية. (تشرح تلك النظرية لماذا وكيف ينصرف الشخص عن متابعة مهمة مُخطَّط لها عندما يلاحظ شيئًا في البيئة لا علاقة له بالمهمة). كذلك ورد في الفصل الخامس جهاز حسيّ حركي يجمع بين برمجة الروبوتات «الموضعية» والشبكات العصبية والبرمجة التطورية. (يساعد هذا الجهاز الروبوت على إيجاد طريقه إلى «المنزل» باستخدام مثلث من الورق المقوّى باعتباره علامة بارزة).

بالإضافة إلى الاستخدامات العملية لتلك النُّهج، فإنها يمكن أن تنير العقل والسلوك والحياة. الشبكات العصبية مفيدة في نمذجة الجوانب العقلية، وفي التعرف على النمط وتعلّمه. يمكن للذكاء الاصطناعي الكلاسيكي (لا سيما عند دمج مع الإحصاءات) أن يضع نموذجًا للتعلّم، وكذلك يمكنه التخطيط والتفكير المنطقي. تُلقِي البرمجة التطورية الضوء على التطور الأحيائي ونمو الدماغ. يمكن استخدام الأتمتة الخلوية والأنظمة الديناميكية في إنشاء نماذج التطور لدى الكائنات الحية. بعض المناهج أقرب إلى علم الأحياء من علم النفس، وبعضها أقرب إلى السلوك التأملي من التفكير التشاوري. لفهم النطاق الكامل للعقل، فسيلزم توافر كل تلك الأنواع وربما أكثر.

لا يهتم كثير من الباحثين في الذكاء الاصطناعي بطريقة عمل العقل؛ حيث إنهم يسعون خلف الكفاءة التكنولوجية وليس الفهم العلمي. حتى وإن كانت أساليب العقل متأصلة في علم النفس، فإنه لا توجد علاقة وثيقة به الآن. لكننا سنرى أن التقدم في الذكاء الاصطناعي للأغراض العامة (الذكاء الاصطناعي العام) يحتاج إلى فهم البنية الحاسوبية للعقل فهمًا عميقًا.

التنبؤ بالذكاء الاصطناعي

تنبأت السيدة آدا لافليس بالذكاء الاصطناعي، تنبأت به السيدة في أربعينيات القرن التاسع عشر. ولمزيد من الدقة، فقد تنبأت بشق من الذكاء الاصطناعي. ركزت السيدة آدا لافليس على الرموز والمنطق، ولم يكن لديها أدنى فكرة بشأن الشبكات العصبية أو

الذكاء الاصطناعي التطوري أو الديناميكي. كذلك لم يكن لديها أي ميل تجاه الهدف النفسي من الذكاء الاصطناعي، بل انصبَّ كل اهتمامها على الجانب التكنولوجي.

قالت — على سبيل المثال — إن الآلة بإمكانها أن «تؤلف مقطوعات موسيقية دقيقة وعلمية مهما كان تعقيدها أو طالت مدتها»، ويمكن أن تعبر أيضًا عن «الحقائق العظيمة للعالم الطبيعي»، وهو ما يُمكن «لحقة مجيدة في تاريخ العلوم». (لذا ما كانت لتتفاجأ حين ترى العلماء بعد قرنين يستخدمون «البيانات الضخمة» وجيلًا برمجية مُبتكرة خصوصًا للتقدم في علم الوراثة والصيدلة وعلم الأوبئة وغيرها ... والقائمة لا تنتهي).

الآلة التي كانت في بالها هي المحرك التحليلي. إنه جهاز مُكوّن من تروس وعجلات مُسنّنة (ولم يكتمل بناؤه قط) من تصميم صديقها المقرّب تشارلز باباج عام ١٨٣٤. وعلى الرغم من أن الجهاز كان مُخصّصًا للجبر والأعداد، فإنه كان معادلًا في الأساس لجهاز كمبيوتر رقمي يُستخدم في الأغراض العامة.

أدركت آدا لافليس احتمالية تعميم المحرك وقدرته على معالجة الرموز التي تمثّل «كل ما في الكون». كذلك وصفت العديد من أساسيات البرمجة الحديثة؛ البرامج المخزنة والإجراءات الفرعية ذات التداخل الهرمي والعنونة والبرمجة الدقيقة والتكرار الحلقي والجمل الشرطية والتعليقات، بل الأخطاء أيضًا. لكنها لم تذكر شيئًا عن كيفية تنفيذ مقطوعة موسيقية، أو تفكير علمي على آلة باباج. نعم كان الذكاء الاصطناعي ممكنًا، ولكن طريقة تحقيقه كانت لا تزال لغزًا.

كيف بدأ الذكاء الاصطناعي

تكشّف اللغز بعد قرن على يد آلان تورينج. في عام ١٩٣٦، أوضح تورينج أن كل عملية حسابية يمكن تنفيذها من حيث المبدأ باستخدام نظام رياضي يُسمّى الآن آلة تورينج العالمية. هذا النظام التخيلي يبني ويعدّل مجموعات من الرموز الثنائية — التي تُمثّل بالرقمين «٠» و«١». بعد فك الشفرة في بلتشلي بارك في أثناء الحرب العالمية الثانية، قضى ما تبقى من أربعينيات القرن العشرين يفكر بشأن كيفية تقريب آلة تورينج التجريدية باستخدام آلة مادية، وكيفية حث تلك الآلة الغريبة للعمل بذكاء (وقد ساعد في تصميم أول جهاز كمبيوتر حديث، واكتمل بمانشستر عام ١٩٤٨).

وعلى خلاف آدا لافليس، قبل تورينج هُدِّي الذكاء الاصطناعي. أراد أن تنفَّذ الآلات الجديدة أشياء مفيدة يُقال عادةً إنها تتطلب الذكاء (ربما استخدام تقنيات غير طبيعية بدرجة كبيرة)، وكذلك تضع نماذج للعمليات التي تحدث في العقل البيولوجي.

نشر ورقة بحثية عام ١٩٥٠، وقال فيها مازحًا: إن اختبار تورينج (انظر الفصل السادس) كان في المقام الأول بيانًا عن الذكاء الاصطناعي. (كُتِب الإصدار الأكمل بعد الحرب بفترة وجيزة، ولكن مُنِع نشره بموجب قانون الأسرار الرسمية). لقد حدَّد الأسئلة الأساسية عن معالجة المعلومات الداخلة في الذكاء (ممارسة الألعاب والإدراك واللغة والتعلم)، وهو ما أعطى تلميحات مُحيرة بشأن ما تحقِّق بالفعل. («تلميحات» فقط لأن العمل في بلتشي بارك كان لا يزال سرّيًا للغاية). اقترحت الورقة مناهج حوسبة — مثل الشبكات العصبية والحوسبة التطورية — ولم تبرز إلا بعد وقت طويل. لكن لم يكن اللغز قد انقشع بعد. كانت تلك ملاحظات عامة إلى حدٍّ كبير؛ ملاحظات برمجية، وليست برامج.

تعزَّز اقتناع تورينج بأن الذكاء الاصطناعي لا بد أن يكون ممكنًا بطريقة أو بأخرى في أوائل أربعينيات القرن العشرين على يد عالم الأعصاب/الطبيب النفسي وارن ماكولو وعالم الرياضيات وولتر بيتس. وفي ورقة بحثية بعنوان «حساب التفاضل والتكامل المنطقي للأفكار الكامنة في النشاط العصبي»، وحدَّ عمل تورينج مع عنصرين آخرين مثيرين للاهتمام (ويعود كلاهما إلى أوائل القرن العشرين)، وهما: منطق القضايا لبرتراند راسل، ونظرية التشابكات العصبية لتشارلز شرينجتون.

الفكرة الأساسية في منطق القضايا هي الثنائية. وفيه يُفترض أن كل جملة (وتُسمَّى أيضًا «افتراضًا») إمَّا صحيحة وإمَّا خطأ. لا يوجد طريق وسط، ولا يُعترف بالشك أو الاحتمال. ولا يُسمح بأكثر من قيمتين من «قيم الحقيقة»، بمعنى أن تكون إمَّا صحيحة وإمَّا خطأ.

إضافة إلى ذلك، تُصاغ الافتراضات المُعقَّدة، وتُنفَّذ الحجج الاستنتاجية باستخدام المعاملات المنطقية (مثل «و» و«أو» و«الجملة الشرطية») حيث يتحدَّد معناها من حيث صواب/خطأ المقترحات المصوغة. على سبيل المثال، إذا رُبط افتراضان (أو أكثر) بحرف العطف «و»، فسيُفترض أن كلا المقترحين صحيحان. إذن، جملة «ماري زوجة توم وفلوسي زوجة بيتر» صحيحة إذا — فقط إذا — كانت عبارة «ماري زوجة توم» وعبارة «فلوسي زوجة بيتر» صحيحتين.

يمكن الجمع بين راسل وشرينجتون إلى جانب ماكولو وبيتس؛ لأن كلا الفريقين وصَف أنظمة ثنائية. حُطِّطت قيم «صح/خطأ» في المنطق إلى نشاط «تشغيل/إيقاف» في

خلايا الدماغ، وإلى «١ / ٠» في الحالات الفردية في آلة تورينج. كان شرينجتون يعتقد أن الخلايا العصبية ليست صارمة بشأن التشغيل/الإيقاف فحسب، ولكن لها حدودًا دنيا ثابتة أيضًا. ومن ثم عُرفت البوابات المنطقية (حوسبة «و» و«أو» و«ليس») بأنها شبكات عصبية يمكن أن تكون مترابطة للتعبير عن مقترحات بالغة التعقيد. أي شيء يمكن ذكره بمنطق القضايا يمكن حوسبته بشبكة عصبية وبآلة تورينج.

باختصار، جُمع علم وظائف الأعصاب والمنطق والحوسبة بعضهم مع بعض، ثم أُضيف إليها علم النفس. اعتقد ماركول وبيتس (كما اعتقد العديد من الفلاسفة حينذاك) أن اللغة الطبيعية أساسها المنطق من حيث الجوهر. لذا كل نتاجهم من التفكير والآراء — بدايةً من الحجة العلمية وحتى الأوهام الفصامية — كان نتاجًا لمحنة نظرياتهم. وبالنسبة إلى علم النفس بأكمله، فقد توقّعوا وقتًا «ستُساهم فيه مواصفات الشبكة [العصبية] في كل ما يمكن تحقيقه في هذا المجال».

كان المعنى الضمني الأساسي واضحًا، وهو «إمكانية تطبيق المنهج النظري نفسه — أي حوسبة تورينج — على الذكاء البشري وذكاء الآلة».

بالطبع اتفق تورينج مع ذلك. لكنه لم يستطع التقدم أكثر بالذكاء الاصطناعي؛ فالتكنولوجيا المتاحة حينذاك كانت بدائية للغاية. لكن في أواسط خمسينيات القرن العشرين، ابتُكرت أجهزة أقوى وأسهل في الاستخدام. المقصود بـ «أسهل في الاستخدام» هنا الأسهل في تحديد أجهزة افتراضية جديدة (مثل لغات البرمجة) يمكن استخدامها بسهولة أكبر على الأجهزة الافتراضية الأعلى مستوى (مثل البرامج التي تحل المسائل الرياضية أو تضطلع بالتخطيط).

بدأت أبحاث الذكاء الاصطناعي الرمزي — إذ بُنيت على بيان تورينج إلى حد كبير — في كلٍّ من أوروبا وأمريكا. ومن العلامات البارزة في أواخر خمسينيات القرن العشرين لاعب الداما الذي ابتكره آرثر صمويل وتصدرت عناوين الصحف، أن اللاعب تعلّم حتى إلحاق الهزيمة بصمويل نفسه. كانت تلك أمانة على أن أجهزة الكمبيوتر يمكن أن تكتسب ذكاءً بشريًا فائقًا يومًا ما، وتتفوق على قدرات مُبرمجها.

ظهرت أمانة أخرى في أواخر خمسينيات القرن العشرين حينما لم تكتفِ آلة النظريات المنطقية أن تثبت ١٨ نظرية من نظريات راسل المنطقية الأساسية فحسب، بل توصّلت إلى دليل مُعتَبَر على إحدى تلك النظريات. كان هذا رائعًا حقًا. وإن كان صمويل لاعب داما متوسط المستوى، فإن راسل كان من رُوَاد العالم في علم المنطق. (سُر راسل

نفسه بهذا الإنجاز، ولكن رفضت مجلة «جورنال أوف سيمبوليك لوجيك» أن تنشر ورقة بحثية بها برنامج كمبيوتر مسمّى على اسم المؤلف، لا سيما أنه لم يبرهن على نظرية جديدة).

سرعان ما تفوّقت آلة حل المسائل العامة على آلة النظرية المنطقية، والتفوق هنا لا يعني أن آلة حل المسائل العامة يمكن أن تتفوق على أصحاب الذكاء الحاد، ولكن تعني أنها ليست محصورة في مجال واحد. وحسب دلالة الاسم، يمكن تطبيق آلة حل المسائل العامة على أي مسألة يمكن تمثيلها (كما هو موضح في الفصل الثاني) من حيث الأهداف الرئيسية والأهداف الفرعية والإجراءات والمشغلون. وكُل إلى المبرمجين تحديد الأهداف والإجراءات والمعاملات ذات الصلة بأي مجال بعينه. ولكن بمجرد إنجاز تلك المرحلة، فإنه يمكن ترك «التفكير المنطقي» للبرنامج.

على سبيل المثال، تمكّنت آلة حل المسائل العامة من حل مسألة «المبشرين وأكلي لحوم البشر». (ثلاثة مبشرين وثلاثة من أكلي لحوم البشر على ضفة نهر، والزورق يسع فردين، فكيف يعبر الجميع النهر من دون أن يتفوق أكلو لحوم البشر في العدد على المبشرين؟) تلك المسألة صعبة حتى على البشر؛ لأنها تتطلب عودة أحدهم كي يعبر النهر. (حاول حلها باستخدام البنسات!)

كانت آلة النظرية المنطقية وآلة حل المسائل العامة أمثلةً أولية على الذكاء الاصطناعي التقليدي الجميل. لا شك أن الزمن عفا عليهما الآن. ولكن كانت الآلتان «جيدتين» أيضًا؛ حيث إنهما تصدّرا استخدام الاستدلال والتخطيط، وكلاهما له أهمية كبيرة في الذكاء الاصطناعي اليوم (انظر الفصل الثاني).

لم يكن الذكاء الاصطناعي التقليدي الجميل هو النوع الوحيد المستلهم من الورقة البحثية التي تحمل اسم «حساب التفاضل والتكامل المنطقي». كانت هذه الورقة البحثية بمثابة تشجيع لمبدأ الترابطية أيضًا. في خمسينيات القرن العشرين، استُخدمت شبكات الخلايا العصبية المنطقية مأكولو-بيتس، سواء المُصمّمة حسب الغرض أو المحاكاة على أجهزة كمبيوتر رقمية (مثل ألبرت أتلي) لوضع نموذج التعليم الترابطي وردود الأفعال المُكيّفة. (وعلى خلاف الشبكات العصبية اليوم، فتلك الشبكات كانت تُجري عمليات المعالجة المحلية دون الموزعة؛ انظر الفصل الرابع).

ولكن النماذج الأولى للشبكة لم يُهيمن عليها منطق الخلية العصبية. فالأنظمة التي نفّذها ريموند بيرل في أواسط خمسينيات القرن العشرين (في أجهزة الكمبيوتر التماثلية)

كانت مختلفة كثيرًا. بدلاً من شبكات البوابات المنطقية المصممة بعناية، بدأ من المجموعات المرتبة الثنائية الأبعاد للوحدات المتصلة عشوائيًا وذات الحدود الدنيا المتفاوتة. كما رأى التنظيم الذاتي العصبي ناتجًا عن موجات التنشيط الديناميكية؛ البناء والنشر والاستمرار والانتها، وأحيانًا التفاعل.

بناءً على ما أدركه بيرل، فإن القول بأن العمليات النفسية التي يمكن نمذجتها بآلة النقاشات السفسطائية لا يعني أن الدماغ يعمل مثل تلك الآلة في الواقع. أشار كلٌّ من ماكولو وبيتس إلى ذلك من قبل. وبعد أربعة أعوام من نشر ورقتهما البحثية الرائدة، نشرنا ورقة بحثية أخرى يقولان فيها إن الديناميكا الحرارية أقرب إلى عمل الدماغ من المنطق. أفسح المنطق الطريق للإحصاء، والوحدات الفردية للوحدات الجماعية، والصفاء الحتمي للتشويش الاحتمالي.

بعبارة أخرى، تحدّثوا عمّا يُطلَق عليه الآن الحوسبة الموزعة التي تحتل الخطأ (انظر الفصل الرابع). رأوا هذا النهج الجديد على أنه «امتداد» لنهجهم السابق وليس مناقضًا له. ولكنه كان أكثر واقعية من الناحية البيولوجية.

السرانية

ذهب تأثير ماكولو في الذكاء الاصطناعي المبكر إلى ما هو أبعد من الذكاء الاصطناعي التقليدي الجميل والترابطية. جعلته معرفته بعلم الأعصاب وكذلك المنطق رائدًا ملهمًا في بناء حركة علم السرانية في أربعينيات القرن العشرين.

ركّز اختصاصيو السرانية على التنظيم الذاتي البيولوجي. وهذا يشمل عدة أنواع من التكيف والتمثيل الغذائي، ومنها التفكير المستقل والسلوك الحركي، وكذلك التنظيم الفسيولوجي (العصبي). تمحور مفهوم تلك الأنواع حول مبدأ «التعليل الدائري» أو الملاحظات. وكان الشاغل الرئيسي هو علم الغائية أو تحديد الأهداف. كانت تلك الأفكار وثيقة الصلة بعضها ببعض، حيث اعتمدت الملاحظات على أوجه الاختلاف في الأهداف؛ بمعنى أن المسافة الحالية عن الهدف كانت تُستخدم لتوجيه الخطوة التالية.

أطلق نوربرت وينر (الذي صمم الصواريخ المضادة للصواريخ بالستية في أثناء الحرب) اسمًا على الحركة عام ١٩٤٨، وعرفها بأنها «دراسة التحكم والاتصال لدى الحيوان والآلة». غالبًا ما استوحى اختصاصيو السرانية الذين أنشؤوا نماذج لأجهزة الكمبيوتر الإلهام من هندسة التحكم وأجهزة الكمبيوتر التماثلية دون المنطق أو الحوسبة

الرقمية. ولكن لم يكن الفرق واضحاً تمام الوضوح. على سبيل المثال، استُخدمت أوجه الاختلاف في الأهداف للتحكم في الصواريخ الموجهة ولتوجيه حل المسائل الرمزية. إضافةً إلى ذلك، استخدم تورينج — بطل الذكاء الاصطناعي الكلاسيكي — المعادلات الديناميكية (التي تصف الانتشار الكيميائي) لتحديد الأنظمة الذاتية التنظيم، التي يمكن أن تظهر فيها بنية جديدة مثل البقع أو التجزئة من أصل متجانس (انظر الفصل الخامس).

من الأعضاء الأوائل الآخرين في الحركة عالم النفس التجريبي كينيث كريك، وعالم الرياضيات جون فون نيومان، وعالم الأعصاب ويليام جراي وولتر وويليام روس أشبي، والمهندس أوليفر سيلفريدج، والطبيب النفسي وعالم الأنثروبولوجيا جريجوري باتسون، وعالم الكيمياء وعالم النفس جوردن باسك.

توفي كريك في حادث ركوب دراجة (عن عمر ٣١ عاماً) عام ١٩٤٣ قبل اختراع أجهزة الكمبيوتر الرقمية، وقد أشار إلى الحوسبة التماثلية لما فكّر في الجهاز العصبي. وبوجه عام، وصف الإدراك والفعل الحركي والذكاء بأنها ملاحظات من «نماذج» في الدماغ. ومفهوم النماذج العقلية — أو التمثيلات — ستكون عامل تأثير كبير في الذكاء الاصطناعي.

وقع فون نيومان في حيرة بشأن التنظيم الذاتي في ثلاثينيات القرن العشرين، وسرّ كثيراً بأول ورقة بحثية كتبها ماركوس وبيتس. وإلى جانب تغيير تصميم جهازه الكمبيوتر الأساسي من التصميم العشري إلى التصميم الثنائي، فقد كيف أفكارهما لشرح التطور والتكاثر البيولوجي. حدّد بعض أنظمة الأتمتة الخلوية، وهي أنظمة تتألف من عدة وحدات حاسوبية، وتغييرات تلك الوحدات تتبع قواعد بسيطة تعتمد على الحالة الحالية للوحدات المجاورة. وبعض تلك الوحدات يمكن أن تنسخ وحدات أخرى. حتى إنه حدّد وحدة نسخ عامة قادرة على نسخ أي شيء بما في ذلك نفسها. وقال إن الأخطاء في النسخ يمكن أن تؤدي إلى التطور.

الأتمتة الخلوية حدّدها فون نيومان بمصطلحات معلوماتية مجردة. ولكن يمكن تجسيدها بعدة طرق مثل الروبوتات الذاتية التجميع، أو الانتشار الكيميائي لتورينج، أو الموجات الفيزيائية لبيرل، أو الحمض النووي الذي سرعان ما أصبح واضحاً.

وبدائيةً من أواخر أربعينيات القرن العشرين، طوّر أشبي جهاز هوميوستات، وهو نموذج كهروكيميائي لآلية الاستتباب الفسيولوجية. ذلك الجهاز المثير للاهتمام يمكن أن يحافظ على حالة توازن تام، بغض النظر عن القيم التي عُيّنَت في البداية لمعلماته البالغ عددها ١٠٠ معلمة (مما يُتيح قرابة ٤٠٠ ألف حالة بدء مختلفة). إنه يوضّح نظرية أشبي

في التكيف الديناميكي، سواء داخل الجسم (لا سيما الدماغ) وبين الجسم والبيئة الخارجية باستخدام طريقة التعلم عن طريق المحاولة والخطأ والسلوك التكيفي.

كان جراي وولتر أيضًا يدرس السلوك التكيفي، ولكن بطريقة مختلفة تمامًا. لقد بنى روبوتات صغيرة تشبه السلاحف، وقد حذت داراتها الحسية الحركية حذو نظرية شرينجتون عن ردود الفعل العصبية. أظهرت تلك الروبوتات الرائدة سلوكيات واقعية مثل البحث عن الضوء وتجنب العقبات والتعليم الترابطي عبر ردود الفعل المتكيفة. وقد عُرضت على عامة الجمهور في مهرجان بريطانيا عام ١٩٥١.

بعد ١٠ أعوام، استخدم سلفريدج (حفيد مؤسس متجر لندن المتعدد الأقسام) طرقًا رمزية لتنفيذ نظام معالجة متوازية في الأساس، ويُسمى بانديمونيوم.

تعلم برنامج الذكاء الاصطناعي التقليدي الجميل هذا التعرف على الأنماط عن طريق توافر العديد من «البرامج الخفية» ذات المستوى الأدنى، ولا يتوقف كل برنامج عن البحث عن مُدخل إدراكي بسيط تنتقل نتائجه إلى برامج خفية ذات مستوى أعلى. وازنت تلك البرامج بين الميزات التي تم التعرف عليها بهدف تحقيق الاتساق (مثل وجود شريطين أفقيين فقط في حرف F)، وهو ما يقلل من أهمية أي ميزات غير مناسبة. قد تتفاوت مستويات الثقة، وهذه مسألة مهمة؛ فالبرامج الخفية ذات النشاط الأعلى لها أكبر تأثير. وفي النهاية، اختار البرنامج الخفي المهيمن أكثر نمط معقول بناءً على الأدلة المتاحة (المتضاربة في كثير من الأحيان). وسرعان ما أثر هذا البحث في كلٍّ من الذكاء الاصطناعي الترابطي والرمزي. (ومن أحدث الفروع نموذج تعلم وكيل التوزيع الذكي LIDA الخاص بالوعي؛ انظر الفصل السادس).

لم يهتم باتسون بالآلات كثيرًا، ولكنه أسس نظرياته في ستينيات القرن العشرين عن الثقافة ومعاقرة الخمر والفصام المزدوج الاتصال على الأفكار المتعلقة بالتواصل (مثل التغذية الراجعة)، التي اختيرت في وقت سابق في الاجتماعات بشأن السبرانية. ومنذ أواسط خمسينيات القرن العشرين، استخدم باسك — الذي وصفه ماركول بأنه «عبقري الأنظمة الذاتية التنظيم» — الأفكار السبرانية والرمزية في العديد من المشروعات المختلفة. تضمنت تلك الأفكار المسرح التفاعلي، والتواصل بين الروبوتات الموسيقية، والمعمارية التي تعلمت وتكيفت مع أهداف مُستخدميها والمفاهيم الذاتية التنظيم كيميائيًا، وتعلم الآلة. تعلم الآلة مكّن الناس من أن يسلكوا طرقًا مختلفة عبر تمثيل معرفي مُعقد؛ ومن ثم كانت مناسبة لكلٍّ من الأنماط الإدراكية التدريجية والشاملة (وتفاوت السماح بدرجات انعدام الصلة) من جانب المتعلم.

باختصار، بُحثت كل أنواع الذكاء الاصطناعي الأساسية، حتى إنها نُفذت بحلول أواخر ستينيات القرن العشرين، وبعضها كان في وقت سابق لذلك بكثير. في الوقت الراهن، يحظى معظم الباحثين المعنيين بتقدير كبير. ولكن ظلَّ طيف تورينج حاضراً في وليمة الذكاء الاصطناعي. وعلى مدار عدة سنين، لم يَكُن الباكون يتذكرهم سوى قلة من مجتمع الباحثين. كاد يُنسى جراي وولتر وأشبهي على وجه الخصوص حتى أواخر ثمانينيات القرن العشرين حين أُشيد بهما (مع تورينج) باعتبارهما أبوي الحياة الاصطناعية. لفهم السبب، يجب معرفة كيف تفرَّق مصمِّمو أجهزة الكمبيوتر.

كيف انقسم الذكاء الاصطناعي

قبل ستينيات القرن العشرين، لم يَكُن ثمة فرق واضح بين أصحاب نماذج اللغة أو التفكير المنطقي وأصحاب نماذج السلوك الحركي المقصود/التكيفي. ولكن البعض عمل في المجالين. (حتى إن دونالد ماكاي اقترح بناء أجهزة كمبيوتر هجينة تجمع بين الشبكات العصبية والمعالجة الرمزية). أيد الجميع بعضهم بعضاً. رأى الباحثون الذين يدرسون التنظيم الذاتي الفسيولوجي أنفسهم أنهم مُنخرطون في المشروع الكلي نفسه مثل زملائهم أصحاب التوجُّه النفسي. وجميعهم حضر اللقاءات نفسها، مثل الندوات العلمية المتعدِّدة التخصصات التي كان يعقدها ماسي في الولايات المتحدة (وقد ترأسها ماكولو في الفترة من ١٩٤٦ حتى ١٩٥١) ومؤتمر لندن الرائد بشأن «ميكنة عمليات التفكير» (الذي نظَّمه أتلي عام ١٩٥٨).

لكن منذ حوالي ١٩٦٠، نشأ انقسام فكري. بوجه عام، ظلَّ المهتمون بـ «الحياة» في الحوسبة السبرانية، وتحوَّل المهتمون بـ «العقل» إلى الحوسبة الرمزية. وبالطبع اهتم المتحمِّسون للشبكات بكلِّ من الدماغ والعقل. ولكنهم درسوا التعليم الترابطي بوجه عام، وليس المحتوى الدلالي أو التفكير المنطقي على وجه التحديد؛ ومن ثَم وقعوا ضمن مجال السبرانية وليس الذكاء الاصطناعي الرمزي. ولسوء الحظ، لم يَكُن هناك قدر كبير من الاحترام المتبادل بين المجموعتين الفرعيتين اللتين يزيد الفارق بينهما.

عندئذٍ، ما كان هناك مناص من ظهور مجموعات اجتماعية مميزة. وكذلك اختلفت الأسئلة النظرية التي كانت تُطرح — سواء البيولوجية (ذات الأنواع المتفاوتة) والنفسية (ذات الأنواع المتفاوتة أيضاً). ولذا أُدخلت المهارات الفنية أيضاً؛ وهي بوجه عام المقارنة

بين المعادلات المنطقية والمعادلات التفاضلية. نمو التخصص زاد من صعوبة التواصل كما زاد من عدم جدواه. وأصبحت المؤتمرات ذات درجة الانتقاء العالية شيئاً من الماضي. وعلى الرغم من ذلك، لم يُكن التقسيم بحاجة إلى أن يكون بهذا السوء. بدأ الشعور السيئ من جانب أصحاب المبدأ السبراني/الترابطي في صورة مزيج من الغيرة المهنية والسخط المحمود. كان الدافع وراء هذا الشعور هو النجاح الأولي الذي حقّقه الحوسبة الرمزية، والاهتمام الصحفي بتصدير المصطلح المستفز «الذكاء الاصطناعي» (الذي صاغه جون مكارثي عام ١٩٥٦ لتسمية ما كان يُطلق عليه قبل ذلك اسم «محاكاة الكمبيوتر») والعجرفة — والضجة غير الواقعية — التي عبّر عنها بعض أصحاب الحوسبة الرمزية. كان معسكر الرمزية أقلّ عداءً في البداية؛ لأنهم كانوا يرون أنفسهم يكسبون منافسة الذكاء الاصطناعي. في الحقيقة، لقد تجاهلوا إلى حد كبير الأبحاث الأولى عن الشبكات على الرغم من أن بعض قادتهم (مثل مارفن مينسكي) بدأ من هذا المجال. لكن في عام ١٩٥٨، طُرحت نظرية طموحة عن الديناميكا العصبية على يد فرانك روزنبلات، وقد نفّذها جزئياً في جهاز بيرسيبترون الكهروضوئي، وتقول النظرية إن أنظمة المعالجة المتوازية قادرة على التعلم الذاتي التنظيم النابع من أساس عشوائي (كما أنه يحتمل الخطأ في البداية). وعلى خلاف نظام بانديمونوم، لا يحتاج هذا النظام أن يحلّ المبرمجون أنماط المدخلات بشكل مسبق. فهذا الشكل الجديد من الترابطية لا يمكن أن يتجاهله فريق الرمزية. ولكن سرعان ما رُفض بازدرء. وحسبما ورد في الفصل الرابع، أطلق مينسكي (مع سيمور بابيرت) نقداً لاذعاً في ستينيات القرن العشرين زعمًا أن أجهزة بيرسيبترون غير قادرة على حساب بعض الأشياء الأساسية. بناءً على ذلك، توقّف تمويل البحوث الخاصة بالشبكات العصبية. تلك النتيجة التي قصدها المهاجمان عن عمد عمّقت العداء بين فرق الذكاء الاصطناعي.

من وجهة نظر العامة، يبدو الآن أن الذكاء الاصطناعي الكلاسيكي كان الخيار الوحيد المتاح حينذاك. ولا أحد يُنكر أن روبوتات السلاحف التي صمّمها جراي وولتر حظيت باستحسان كبير في مهرجان بريطانيا. أحدثت الصحافة ضجة حول جهاز بيرسيبترون الذي صمّمه روزنبلات في أواخر خمسينيات القرن العشرين، كما أحدثت حول نظام أدالين الذي صمّمه برنارد ويدرو للتعلم بالأنماط (بناءً على معالجة الإشارات). ولكن فريق الرمزية أحمَد ذلك الاهتمام تماماً. ومن ثمّ تصدر الذكاء الاصطناعي الرمزي قنوات الإعلام في ستينيات وسبعينيات القرن العشرين (كما أثر في فلسفة العقل كذلك).

لم يستمر الموقف. ظهرت الشبكات العصبية على الساحة مرةً أخرى عام ١٩٨٦ (انظر الفصل الرابع)، مثل أنظمة معالج البيانات المبرمج (التي تُجري المعالجة الموزعة المتوازية). يعتقد كثير من غير المختصين — وبعض المختصين الذين من المفترض أنهم على معرفة أفضل — أن هذا النهج جديد تمامًا. ومن ثم استهوى الخريجين، وجذب انتباه عدد كبير من أهل الصحافة (والفلسفة). والآن، أصحاب الذكاء الاصطناعي الرمزي هم من انزعجوا. كان معالج البيانات المبرمج هو السائد، وشاع أن الذكاء الاصطناعي الكلاسيكي قد فشل.

بالنسبة إلى بقية أنصار السبرانية، فقد انخرطوا مرةً أخرى في المجال بما أسَمَوْه الحياة الاصطناعية عام ١٩٨٧. وتبعهم الصحفيون والخريجون. ومن ثم وقع الذكاء الاصطناعي الكلاسيكي في مأزق مرةً أخرى.

لكن في القرن الحادي والعشرين، أصبح واضحًا أن تنوع الأسئلة يحتاج إلى تنوع الإجابات — لكل مقام مقال. وعلى الرغم من بقاء آثار العداء القديم، توجد فسحة للاحترام — وحتى التعاون — بين النهج المختلفة. على سبيل المثال، يُستخدم «التعلم العميق» في بعض الأحيان في الأنظمة القوية التي تجمع بين المنطق الرمزي والشبكات الاحتمالية المتعددة الطبقات، وغيرها من النهج المختلطة تتضمن نماذج طموحة من الوعي (انظر الفصل السادس).

بناءً على التنوع الثري للأجهزة الافتراضية التي تشكّل العقل البشري، فلا ينبغي أن يندهش المرء كثيرًا.

الفصل الثاني

كأن الذكاء العام هو الكأس المقدسة

ينطوي الذكاء الاصطناعي الحديث على العديد من الميزات. فهو يوفر كثيرًا من الأجهزة الافتراضية التي تُجري العديد من أنواع معالجة المعلومات. لا يوجد سر رئيسي هنا، ولا توجد تقنية أساسية توحد المجال؛ أي إن العاملين في الذكاء الاصطناعي يعملون في مجالات كثيرة التنوع، ولا يتشاركون سوى القليل من حيث الأهداف والطرق. ولا يذكر هذا الكتاب سوى نزر يسير من التطورات الحديثة. باختصار، نطاق مناهج الذكاء الاصطناعي واسع للغاية.

يمكن القول إن الذكاء الاصطناعي حقق نجاحات باهرة. فنطاق استخداماته أيضًا واسع للغاية. توجد مجموعة من تطبيقات الذكاء الاصطناعي المُصمَّمة لأداء عدد لا حصر له من المهام، وتلك التطبيقات يستخدمها الإنسان العادي والمحترف على حدٍّ سواء في كل مناحي الحياة. والعديد منها يتفوق على البشر حتى أكثرهم خبرة. ومن هذا المنطلق، كان التقدم الذي حققه الذكاء الاصطناعي رائعًا.

لكن لم يستهدف رواد الذكاء الاصطناعي الأنظمة المتخصصة فحسب. فقد كانوا يتطلعون إلى إنشاء أنظمة تمتاز بذكاء عام. وستواجه كل نماذج القدرات الشبيهة بقدرات الإنسان — الرؤية والتفكير المنطقي واللغة والتعلم وما إلى ذلك — مجموعة كاملة من التحديات. إضافة إلى ذلك، تتكامل تلك القدرات عندما تكون ملائمة.

بالحكم بناءً على تلك المعايير، فإن التقدم سيُثير الإعجاب بدرجة أقل بكثير. أدرك جون مكارثي حاجة الذكاء الاصطناعي إلى «منطق عقلائي» في وقت مبكر للغاية. وقد تحدث عن «الشمولية في الذكاء الاصطناعي» في الخطابين اللذين ألقاهما حين نال جائزة تورينج المرموقة في عامي ١٩٧١ و١٩٨٧، لكن كان محتوى الخطابين شكوى وليس احتفالاً. وحتى عام ٢٠١٨، لم تُزل أسباب شكواه بعد.

يشهد القرن الحادي والعشرون عودة الاهتمام بالذكاء الاصطناعي العام؛ إذ تقوده التطورات الأخيرة في قدرات أجهزة الكمبيوتر. وإذا تحقّق ذلك، فسيقلّ اعتماد أنظمة الذكاء الاصطناعي على حيل البرمجة ذات الأغراض الخاصة، بل إنها ستستفيد من القدرات العامة للتفكير المنطقي والإدراك، بالإضافة إلى اللغة والإبداع والعاطفة (وكل ذلك تناولناه في الفصل الثالث).

لكن القول أسهل من الفعل. لا يزال الذكاء العام تحدّيًا كبيرًا، ولا يزال بعيد المنال. الذكاء الاصطناعي العام هو الكأس المقدسة لذلك المجال.

أجهزة الكمبيوتر العملاقة ليست كافية

لا شك أن أجهزة الكمبيوتر العملاقة في الوقت الراهن تمثّل يد العون لكل من يريد تحقيق هذا الحلم. الانفجار التوافقي يُطلّب فيه عدد عمليات حسابية أكثر مما تُنفَّذ بالفعل، ولكنه لم يُعدّ يمثل تهديدًا ثابتًا كما كان. ومع ذلك، المسائل لا تُحلّ دومًا بمجرد رفع قدرات الكمبيوتر.

كثيرًا ما يلزم توفير طرق جديدة لحل المسائل. إضافة إلى ذلك، إذا كانت هناك طريقة معيّنة نجاحها مضمون من حيث المبدأ، فقد تحتاج قدرًا كبيرًا من الوقت والذاكرة أو أحدهما كي تنجح عمليًا. وأوردنا ثلاثة أمثلة (بشأن الشبكات العصبية) في الفصل الرابع.

الكفاءة مهمة أيضًا؛ فكلما قلّ عدد عمليات الحوسبة، كان ذلك أفضل. باختصار، يجب جعل المسائل سهلة الحل.

هناك عدة استراتيجيات لذلك. احتلّت كل تلك الاستراتيجيات مركز الريادة في الذكاء الاصطناعي الرمزي الكلاسيكي — أو الذكاء الاصطناعي التقليدي الجميل — ولا تزال جميعها ضرورية في الوقت الراهن.

إحدى تلك الاستراتيجيات أن نوجّه الانتباه إلى جزء واحد من «مساحة البحث» (تمثيل الكمبيوتر الخاص بالمسألة، وهو الذي من المفترض أن الحل موجود بداخله). استراتيجية أخرى، وهي أن نصغّر مساحة البحث عن طريق تبسيط الافتراضات. استراتيجية ثالثة، وهي أن نرتّب البحث ترتيبًا فعليًا. استراتيجية رابعة، وهي إنشاء مساحة بحث مختلفة عن طريق تمثيل المسألة بطريقة جديدة.

تتضمن تلك النُّهج «الاستدلال والتخطيط والتبسيط الرياضي وتمثيل المعرفة» على التوالي. تتناول الأقسام الخمسة الآتية تلك الاستراتيجيات العامة للذكاء الاصطناعي.

البحث الاستدلالي

لفظة Heuristic (الاستدلال) تنحدر من أصل واحد، مثل لفظة Eureka! (وجدتها!) إنها تنحدر من اللغة اليونانية، وتعني «وجد» أو «اكتشف». برز الاستدلال في مجال الذكاء الاصطناعي التقليدي الجميل المبكر، وكثيراً ما يُعتقد أنه عبارة عن «جِل برمجية». ولكن لم يتأصل المصطلح مع البرمجة، بل إنه معروف منذ وقت طويل بين أوساط علماء المنطق وعلماء الرياضيات.

الاستدلال يسهّل حل المسألة، سواء مع البشر أو الآلات. وفي الذكاء الاصطناعي، يسهل حل المسألة بتوجيه البرنامج صوب أجزاء معيّنة من مساحة البحث وإبعاده عن أجزاء أخرى.

العديد من عمليات الاستدلال — بما في ذلك العمليات التي استُخدمت في أوائل اختراع الذكاء الاصطناعي — عبارة عن قواعد عامة؛ ومن ثم لا يُضَمّن نجاحها. ربما يكمن الحل في جزء من مساحة البحث، وقد وجّه الاستدلال النظام إلى أن يتجاهل ذلك الجزء. على سبيل المثال، «حماية الملكة» قاعدة مفيدة للغاية في الشطرنج، ولكن ينبغي مخالفتها في بعض الأحيان.

ربما تثبت ملاءمة قواعد أخرى من حيث المنطق أو الرياضيات. يهدف قدرٌ كبير من العمل في الذكاء الاصطناعي وعلوم الكمبيوتر اليوم إلى تحديد خصائص يمكن إثبات صحتها في البرامج. هذا جانب من «الذكاء الاصطناعي الصديق»؛ لأنه يُمكن أن تتعرّض سلامة البشر إلى الخطر إن استُخدمت أنظمة غير موثوقة منطقياً (انظر الفصل السابع). سواء كان النظام موثقاً أم لا، فالاستدلال جانب ضروري في البحوث الخاصة بالذكاء الاصطناعي. تعتمد زيادة التخصصية في الذكاء الاصطناعي المذكورة فيما سبق اعتماداً جزئياً على تعريف الاستدلال الجديد الذي يُمكن أن يحسّن الكفاءة بدرجة مذهلة، ولكن لا يحدث هذا إلا في مسألة — أو مساحة بحث — تتسم بدرجة تقييد عالية. الاستدلال ذو النجاح الباهر قد لا يكون مناسباً كي «تقترضه» برامج أخرى تعمل بالذكاء الاصطناعي. واستناداً إلى عدة استدلالات، فقد يكون ترتيب تطبيقها مهماً. على سبيل المثال، يجب أن تؤخذ «حماية الملكة» في الاعتبار قبل «حماية البيدق»، على الرغم من أن هذا الترتيب

قد يؤدي إلى كارثة في بعض الأحيان. عمليات الترتيب المختلفة ستؤدي إلى أشجار بحث مختلفة عَبر مساحة البحث. لذا فإن تحديد الاستدلالات وترتيبها من المهمات البالغة الأهمية في الذكاء الاصطناعي الحديث. (الاستدلال جزء بارز في علم النفس المعرفي أيضًا. يوضح العمل الرائع في «الاستدلال السريع والمقتصد» على سبيل المثال إلى أي مدى زودنا التطور بطرق فعّالة في الاستجابة للبيئة).

الاستدلال يجعل البحث بالقوة عَبر مساحة البحث بالكامل غير ضروري. ولكنه مصحوب في بعض الأحيان بعمليات بحث بالقوة، وإن كانت محدودة. برنامج الشطرنج «ديب بلو» من شركة «آي بي إم» أسعد العالم لما هزم بطل العالم في الشطرنج جاري كاسبروف عام ١٩٩٧، وقد استخدم رقائق أجهزة دقيقة تعالج ٢٠٠ مليون موضع في الثانية؛ ومن ثم يولد كل حركة محتملة للثمانية المقبلة.

ومع ذلك، اضطر إلى استخدام الاستدلال لتحديد «أفضل» حركة من بين الثمانية. وبما أنه لا يمكن الاعتماد على الاستدلال، لم يتمكن حتى برنامج «ديب بلو» من هزيمة كاسبروف في كل مرة.

التخطيط

التخطيط أيضًا جزء بارز في الذكاء الاصطناعي اليوم، كما أنه أصبح واسع الانتشار لا سيما في الأنشطة العسكرية. كانت وزارة الدفاع الأمريكية تموّل الجزء الأكبر من بحوث الذكاء الاصطناعي حتى وقت قريب، وقالت إن الأموال التي وفّرتها (بفضل تخطيط الذكاء الاصطناعي) في لوجستيات المعركة في حرب العراق الأولى فاقت كل الاستثمارات السابقة. لا يقتصر التخطيط على الذكاء الاصطناعي، بل جميعنا نستخدمه. لنضرب مثالاً بحزم الأمتعة للخروج في إجازة. يجب العثور على جميع الأغراض التي تريد أخذها، وربما لن تعثر عليها في مكان واحد. كذلك قد تُضطر إلى شراء بعض الأغراض الجديدة (مثل كريم واقٍ من أشعة الشمس). يجب أن تقرّر هل ستجمع الأغراض كلها في مكان واحد (كأن تجمعها على السرير أو الطاولة)، أم ستضع العنصر الذي تعثر عليه في حقيبة الأمتعة. سيعتمد هذا القرار جزئيًا على ما إذا كنت تريد أن تجعل الملابس آخر الأغراض لمنع تجميدها. ستحتاج إلى حقيبة ظهر أو حقيبة سفر أو ربما حقيبتين؛ كيف ستقرّر؟ كانت تلك الأمثلة المدروسة بوعي في عقل مُبرمجي الذكاء الاصطناعي التقليدي الجميل؛ إذ اتخذوا التخطيط أسلوبًا من أساليب الذكاء الاصطناعي. يرجع السبب في ذلك

إلى أن الرواد المسئولين عن آلة النظرية المنطقية (انظر الفصل الأول) وآلة حل المسائل العامة كانوا مهتمين في الأساس بعلم النفس المعني بالتفكير البشري. لا تعتمد أدوات التخطيط الحديثة ذات الذكاء الاصطناعي اعتمادًا كبيرًا على الأفكار التي يولدها الاستبطان الواعي أو الملاحظة التجريبية. وخططهم معقدة أكثر مما كان محتملاً في بادئ الأمر. ولكن الفكرة الأساسية واحدة.

تحدد الخطة سلسلة من الإجراءات الممثلة على المستوى العام وهدفًا نهائيًا، بالإضافة إلى أهداف فرعية وأهداف متفرعة من الأهداف الفرعية؛ ومن ثم لا تُدرس كافة التفاصيل مرة واحدة. التخطيط في مستوى التجرد الملائم يمكن أن يؤدي إلى تشذيب الشجرة داخل مساحة البحث؛ ومن ثم لا نحتاج بتاتًا إلى دراسة بعض التفاصيل. وفي بعض الأحيان يكون الهدف النهائي نفسه خطة عملية؛ فقد يتمثل في خطط عمليات التسليم إلى المصنع أو أرض المعركة ومنهما. وفي أوقات أخرى، يتمثل في إجابة عن سؤال، مثل التشخيصات الطبية. بالنسبة إلى أي هدف محدد ومواقف متوقعة، يحتاج برنامج التخطيط إلى ما يأتي: قائمة الإجراءات، وهي المعاملات الرمزية أو أنواع الإجراءات (الممثلة باستيفاء المعلومات المشتقة من المسألة)، وكل إجراء قد يحدث بعض التغييرات ذات الصلة؛ ومجموعة المتطلبات الأساسية لكل إجراء (كي تقارن لفهم شيء ما، يجب أن يكون سهل المنال)؛ والاستدلالات لتحديد أولويات التغييرات المطلوبة وترتيب الإجراءات. إذا قرّر البرنامج اتخاذ إجراء مُعَيّن، فقد يُضطر إلى تعيين هدف فرعي جديد كي يلبي المتطلبات الأساسية. يمكن تكرار عملية صياغة الأهداف مرارًا وتكرارًا.

التخطيط يمكن البرنامج — والمستخدم البشري — من اكتشاف الإجراءات التي اتُخذت بالفعل وسبب اتخاذها. يشير «السبب» إلى التسلسل الهرمي للهدف؛ اتُخذ «هذا» الإجراء لتلبية «ذلك» المطلب الأساسي لتحقيق الهدف الفرعي «كذا وكذا». وبوجه عام، توظف أنظمة الذكاء الاصطناعي تقنيات «التسلسل الأمامي» و«التسلسل العكسي» التي تشرح كيف عثر البرنامج على الحل. هذا يساعد المستخدم على أن يحكم هل الإجراء/المشورة من البرنامج ملائم أم لا.

تمتلك بعض أدوات التخطيط الحالية عشرات الآلاف من سطور التعليمات البرمجية؛ إذ تُحدد مساحات البحث الهرمية على عدد هائل من المستويات. وتلك الأنظمة غالبًا ما تكون مختلفة كثيرًا عن أدوات التخطيط الأولى.

على سبيل المثال، معظم تلك الأدوات لا تفترض أنه يمكن العمل على الأهداف الفرعية كلٌّ على حدة (بمعنى أن المسائل قابلة للتجزئة إلى أجزاء فرعية كثيرة). ولكن في الحياة

الواقعية، قد لا تتحقق نتيجة نشاط موجه حسب الهدف بسبب نشاط آخر. أما اليوم، فتستطيع أدوات التخطيط التعامل مع المسائل القابلة للتجزئة جزئياً؛ إنها تعمل على الأهداف الفرعية كلٌّ على حدة، ولكنها تجري معالجة إضافية لدمج الخطط الفرعية الناتجة إذا لزم الأمر.

لا تستطيع أدوات التخطيط الكلاسيكية التعامل إلا مع المسائل التي تكون بيئتها قابلة للرصد بالكامل وحتمية ومحددة وثابتة. ولكن بعض أدوات التخطيط الحديثة يمكن أن تتماشى مع البيئات القابلة للرصد جزئياً (أي إن نموذج النظام للعالم قد يكون غير مكتمل أو غير صحيح أو كليهما) والاحتمالية. في تلك الحالات، يجب أن يرصد النظام الموقف المتغير أثناء التنفيذ، بحيث يجري التغييرات في الخطة — وفي «معتقداتها» بشأن العالم — حسب الحاجة. يمكن لبعض أدوات التخطيط الحديثة أن تفعل ذلك على مدار فترات طويلة، إنها تدخل في عملية متواصلة لتكوين الأهداف وتنفيذها وتعديلها والتخلي عنها تكيّفاً مع البيئة المتغيرة.

أضيفت العديد من عمليات التطوير الأخرى إلى التخطيط الكلاسيكي، ولا تزال تضاف التطويرات. إذن، قد يكون مفاجئاً أن التخطيط رفضه بعض مُصممي الروبوتات رفضاً صريحاً في ثمانينيات القرن العشرين؛ ومن ثمّ تمت التوصية بالروبوتات «الكائنة» بدلاً من ذلك (انظر الفصل الخامس). وكذلك رُفضت فكرة التمثيل الداخلي، مثل رفض الأهداف والإجراءات المحتملة. ولكن كان هذا النقد خاطئاً إلى حدٍّ كبير. غالباً ما تحتاج الروبوتات إلى التخطيط بالإضافة إلى الإجابات التفاعلية الخالصة، مثل بناء روبوتات تلعب كرة القدم.

التبسيط الرياضي

لما كان الاستدلال يترك مساحة البحث كما هي (بأن يجعل البرنامج لا يركّز إلا على جزء منها)، فإن تبسيط الافتراضات يبني مساحة بحثٍ غير واقعية، ولكن يُمكن حلها من منظور الرياضيات.

بعض تلك الافتراضات رياضية. من الأمثلة على ذلك افتراض «المتغيرات المستقلة ومتشابهة التوزيع» المستخدمة في تعلم الآلة. وهذا يمثل الاحتمالات في البيانات وكأنها أبسط بكثير مما هي عليه في الواقع.

تكمن ميزة التبسيط الرياضي عند تحديد مساحة البحث في إمكانية استخدام الطرق الرياضية؛ أي إنه يمكن تحديدها بوضوح ويسهل فهمها على الأقل بالنسبة إلى علماء الرياضيات. ولكن هذا لا يعني أن أي بحث محدّد رياضياً سيكون مفيداً. وكما أشرنا مسبقاً، قد تكون الطريقة المضمونة رياضياً لحل كل مسألة ضمن فئة معيّنة غير قابلة للاستخدام في الحياة الواقعية؛ لأنها قد تحتاج وقتاً غير محدود لحلها. ولكن قد تقترح حلولاً تقريبية عملية أكثر؛ انظر مناقشة «الانتشار العكسي» في الفصل الرابع.

افتراضات التبسيط غير الرياضية في الذكاء الاصطناعي كثيرة للغاية، ولكن غالباً ما يُسكت عنها. ومن الأمثلة على ذلك الافتراض (الضمني) بأنه يمكن تحديد المسائل وحلها من دون أخذ العاطفة في الاعتبار (انظر الفصل الثالث). والعديد من الافتراضات الأخرى داخلة في تمثيل المعرفة العام المستخدم في تحديد المهمة.

تمثيل المعرفة

غالباً ما يكون الجزء الأصعب في حل مسألة الذكاء الاصطناعي هو عرض المسألة على النظام للمرة الأولى. حتى إذا بدا أن أحداً يستطيع التواصل مع برنامج مباشرةً، ربما بالتحدث بالإنجليزية إلى «سيري»، أو بالكتابة بالفرنسية في محرك البحث جوجل، فلا يمكنه التواصل معه. وسواء كان الشخص يتعامل مع نصوص أو مع صور، فإنه يجب عرض المعلومات (المعرفة) المعنية على النظام بطريقة تفهمها الآلة؛ أو بعبارة أخرى، بطريقة تمكّن الآلة من التعامل معها. (وما إذا كان ذلك الفهم حقيقياً أم لا، فقد تناولناه في الفصل السادس).

الطرق التي يفهم بها الذكاء الاصطناعي المعلومات كثيرة ومتنوعة. بعضها عبارة عن تطويرات/تباينات في الطرق العامة لتمثيل المعرفة المقدم في الذكاء الاصطناعي التقليدي الجميل. وهناك طرق أخرى ذات درجة تخصّص عالية، وهي مصمّمة خصوصاً لفئة محدودة من المسائل. على سبيل المثال، قد توجد طريقة جديدة لتمثيل صور الأشعة السينية أو الصور الفوتوغرافية لفئة معيّنة من الخلايا السرطانية؛ ومن ثم فهي تُصمّم بعناية لتوفير طريقة ذات درجة تخصّص عالية في التفسيرات الطبية (ومن ثم لا تصلح للتعرف على صور القطط أو حتى عمليات المسح باستخدام التصوير المقطعي المحوسب).

في البحث عن الذكاء الاصطناعي العام، تمثل الأساليب العامة أهمية قصوى. استندت تلك الأساليب في البداية إلى بحث نفسي عن الإدراك البشري، وتلك الأساليب تتضمن ما

يأتي: مجموعات من القواعد «الشرطية»، وتمثيلات المفاهيم الفردية، وسلاسل الإجراءات النمطية، والشبكات الدلالية، والاستدلال بالمنطق أو الاحتمال. لندرس كل أسلوب على حدة. (هناك شكل آخر من تمثيل المعرفة يُسمّى الشبكات العصبية، وسنتناوله في الفصل الرابع).

البرامج القائمة على القواعد

في البرمجة القائمة على القواعد، تُمثّل مجموعة المعرفة/المعتقد في صورة قواعد «شرطية» تربط فعل الشرط بجواب الشرط: إذا تحقّق هذا الشرط، فاتخذ هذا الإجراء. يستند هذا الشكل من تمثيل المعرفة إلى المنطق الصوري (أنظمة إنتاج إيميل بوست). ولكن رَوّاد الذكاء الاصطناعي ألين نيويل وهيربرت سيمون يعتقدان أن هذا الشكل يمثّل أساس الجانب النفسي لدى الإنسان بوجه عام.

قد يكون كلّ من فعل الشرط وجواب الشرط مُعقّدين؛ وهو ما يؤدي إلى تحديد ارتباط (أو فصل) بعض العناصر أو ربما العديد من العناصر. وإذا تحقّقت بعض أفعال الشرط في آن واحد، فستُعطى الأولوية لأشمل ارتباط. لذا، سيحتل فعل الشرط «إذا كان الهدف طهي لحم بقري مشوي وبودينج يوركشاير» الصدارة على فعل الشرط «إذا كان الهدف طهي لحم بقري مشوي»، كما أن إضافة «وإضافة ثلاثة أنواع من الخضراوات» إلى فعل الشرط سيجعله يعلو.

البرامج القائمة على القواعد لا تحدّد ترتيب الخطوات مقدّمًا. بل إن كل قاعدة تنتظر حتى يحفزها فعل الشرط. ومع ذلك، يمكن استخدام تلك الأنظمة لوضع الخطط. وإذا لم تُستخدم في ذلك، فسينحسر استخدامها في الذكاء الاصطناعي. ولكنها تؤدي تلك المهمة بطريقة مختلفة عن الطريقة القديمة والمعتادة في البرمجة (ويُطلق عليها في بعض الأحيان «التحكم التنفيذي»).

في البرامج ذات التحكم التنفيذي، تُمثّل الخطط صراحةً. يحدّد المبرمج سلسلة من التعليمات الموصلة إلى الهدف حتى تتّبع خطوة بخطوة، وبترتيب زمني صارم: «افعل كذا ثم كذا، ثم انظر هل س صحيحة أم لا، وإذا كانت صحيحة فافعل كذا وكذا، وإذا لم تُكن صحيحة فافعل كذا وكذا».

في بعض الأحيان، تكون «كذا» أو «كذا وكذا» تعليمات صريحة لتعيين هدف أساسي أو هدف فرعي. على سبيل المثال، الروبوت المُعيّن له هدف مغادرة الغرفة يمكن ضبط

تعليماته لتعيين هدف فرعي، وهو فتح الباب؛ ثم إذا نتج عن فحص الحالة الحالية أن الباب مُغلق، فسيتعين هدف مُتفرّع عن الهدف الفرعي، وهو جذب مقبض الباب. (الطفل البشري قد يحتاج إلى هدف مُتفرّع عن الهدف المُتفرّع عن الهدف الفرعي؛ أي يطلب من شخص كبير أن يجذب مقبض الباب الذي لا يستطيع الوصول إليه، وقد يحتاج الرضيع إلى عدة أهداف بمستويات أقل من أجل الوصول إلى هذا الهدف).

يمكن للبرنامج القائم على القواعد أن يفكر أيضاً في طريقة للهروب من الغرفة. لكن قد لا يُمثل التسلسل الهرمي للخطّة في سلسلة خطوات صريحة ذات ترتيب زمني، بل في صورة بنية منطقية «ضمنية» في مجموعة من القواعد «الشرطية» التي يتألف منها النظام. وقد يتطلب فعل الشرط تعيين الهدف كذا بالفعل (إذا أردت أن تفتح الباب، ولست بالطول الكافي). وبالمثل، قد يتضمن جواب الشرط تعيين هدف رئيسي أو هدف فرعي جديد (فاطلب من أحد الكبار). ستُنشّط المستويات الدنيا تلقائياً (إذا أردت أن تطلب من أحد أن يفعل لك شيئاً، فعين هدف الاقتراب منه).

بالطبع، يجب على المبرمج إدراج كل القواعد «الشرطية» ذات الصلة (في المثال الذي معنا، تتعامل القواعد مع الأبواب ومقابض الأبواب). لكنه لا يحتاج إلى توقّع كل التعقيدات المنطقية المحتملة لتلك القواعد. (هذه نقمة ونعمة في الوقت نفسه؛ لأنه قد تظل التناقضات المحتملة غير مكتشفة لوقت طويل).

تُنشر الأهداف الرئيسية/الفرعية النشطة على «لوحة» مركزية يسهل على النظام بأكمله الوصول إليها. لا تتضمن المعلومات الظاهرة على اللوحة الأهداف النشطة فحسب، بل تتضمن المدخلات الإدراكية والأنماط الأخرى للمعالجة الحالية. (أثّرت تلك الفكرة في النظرية العصبية النفسية الرائدة الخاصة بالوعي، وكذلك نموذج الذكاء الاصطناعي للوعي القائم على تلك النظرية؛ انظر الفصل السادس).

شاع استخدام البرامج القائمة على القواعد في «الأنظمة الخبيرة» الرائدة في أوائل سبعينيات القرن العشرين. من تلك الأنظمة «ميسين» (MYCIN)، وكان يقدم المشورة للأطباء لتحديد الأمراض المعدية ووصف المضادات الحيوية؛ ونظام «دندرال» (DENDRAL) الذي يُجري تحليلاً طيفياً للجزيئات داخل مساحة مُعيّنة من الكيمياء العضوية. على سبيل المثال، يُجري نظام «ميسين» التشخيص الطبي عن طريق مطابقة الأعراض والخصائص الجسدية الأساسية (فعل الشرط) للخروج بنتائج التشخيص واقتراح المزيد من الاختبارات أو الأدوية (جواب الشرط). كانت تلك البرامج هي الخطوة

الأولى التي نأى بها الذكاء الاصطناعي عن أمل التعميم واقترب بها إلى التخصيص. كذلك كان البرنامجان الخطوة الأولى نحو حلم آدا لافليس بالعلوم التي تصنعها الآلة (انظر الفصل الأول).

شكل تمثيل المعرفة القائم على القواعد يمكّن من بناء البرامج تدريجيًا؛ حيث إن المبرمج — أو ربما نظام الذكاء الاصطناعي العام نفسه — يتعلّم المزيد عن النطاق. يمكن إضافة قاعدة جديدة في أي وقت. ومن ثم لا يلزم كتابة البرنامج من البداية. لكن هناك مشكلة خفية. إذا لم تتسق القاعدة الجديدة مع القواعد الموجودة بالفعل، فلن يقوم النظام بالإجراء المفترض على الدوام. وربما لا «يقرب» الإجراء المفترض أن يؤديه. وعند التعامل مع مجموعة صغيرة من القواعد، يسهل تجنب تلك التناقضات المنطقية، ولكن الأنظمة الأكبر أقل شفافية.

في سبعينيات القرن العشرين، استُمدت القواعد «الشرطية» الجديدة من المحادثات الجارية مع الخبراء من البشر ممّن طُلب منهم شرح قراراتهم. واليوم، لا يُستمد العديد من القواعد من الاستبطان الواعي. ولكنها تتسم بفاعلية أكبر. تتراوح الأنظمة الخبيرة الحديثة (وهذا المصطلح نادر الاستخدام اليوم) ما بين البرامج المستخدمة في البحوث العلمية والتجارة حتى التطبيقات المتواضعة على الهواتف. يتفوّق العديد من تلك التطبيقات على التطبيقات السابقة؛ لأنها تستفيد من الأشكال الإضافية لتمثيل المعرفة، مثل الإحصاءات وأدوات التعرّف البصرية ذات الأغراض الخاصة واستخدام البيانات الضخمة أو أيّ من ذلك (انظر الفصل الرابع).

بإمكان تلك البرامج أن تساعد الخبراء من البشر أو حتى أن تحل محلهم في مجالات بالغة التعقيد. وفي الوقت الراهن، هناك أمثلة لا حصر لها؛ إذ تُستخدم البرامج لمساعدة المحترفين العاملين في العلوم والطب والقانون ... وحتى تصميم الملابس. (وتلك الأخبار ليست جديدة تمامًا؛ انظر الفصل السابع).

الأطر، متجهات الكلمات، البرامج النصية، الشبكات الدلالية

هناك طرق أخرى شائعة في تمثيل المعرفة تهتم بالمفاهيم الفردية، وليس بالنطاقات بأكملها (مثل التشخيص الطبي أو تصميم الملابس).

على سبيل المثال، يستطيع المرء أن يُخبر الكمبيوتر ما هي الغرفة، بتحديد بنية بيانات تسلسلية (تُسمّى في بعض الأحيان بالإطار). وهذا يمثل الغرفة بأن لها «أرضية

وسقفًا وجدرانًا وأبوابًا ونوافذ وأثاثًا (سريراً وحمامًا ومائدة طعام). «الغرف الحقيقية تتفاوت في عدد الجدران والأبواب والنوافذ؛ ومن ثم ترك «فراغات» في الإطار يتيح كتابة أرقام مُعيَّنة وكذلك تقديم تعيينات افتراضية (أربعة جدران، باب واحد، نافذة واحدة). بنى البيانات هذه يمكن أن يستخدمها الكمبيوتر للبحث عن الأدوات المماثلة أو الرد عن الأسئلة أو المشاركة في محادثة أو كتابة قصة أو فهمها. إنها أساس برنامج سي واي سي (CYC): إنه برنامج طموح — ويقول البعض إنه مُفرط في الطموح — إذ يحاول أن يمثل المعارف البشرية.

لكن الأطر قد تكون مُضلَّة. تنطوي التعيينات الافتراضية على سبيل المثال على إشكاليات. (بعض الغرف ليس فيها نوافذ، والغرفة المكشوفة ليس لها باب). الأسوأ من ذلك، ماذا عن المفاهيم اليومية مثل «السقوط» أو «الانسكاب»؟ يمثل الذكاء الاصطناعي الرمزي معرفتنا المنطقية «للفيزياء المفرطة في البساطة» عن طريق بناء أطر تحوّل تلك الحقائق إلى رموز، مثل حقيقة أن الجسم المادي سيسقط إذا لم يكن مدعومًا. ولكن لن يسقط بالون الهيليوم. والسماح صراحةً لتلك الحالات مهمة لا تتقطع أبدًا.

في بعض التطبيقات التي تستخدم التقنيات الحديثة للتعامل مع البيانات الضخمة، يمكن تمثيل مفهوم مفرد في صورة مجموعة — أو «سحابة» — تتألف من مئات أو آلاف المفاهيم التي تكون مرتبطة في بعض الأحيان، مع تمييز احتمالات العديد من الارتباطات المزدوجة؛ انظر الفصل الثالث. وبالمثل، يمكن تمثيل المفاهيم الآن باستخدام «متجهات الكلمات» بدلًا من الكلمات. وهنا، يكتشف نظام (التعلم العميق) السمات الدلالية التي تساهم في العديد من المفاهيم المختلفة وتربطها بعضها ببعض، كما أن النظام يتنبأ بالكلمة التالية، مثلما يحدث في الترجمة الآلية على سبيل المثال. ومع ذلك، يتعذر استخدام هذه التمثيلات حتى الآن في التفكير المنطقي أو المحادثة، مثل الإطارات الكلاسيكية.

تشير بعض بنى البيانات (وتُسمى البرامج النصية) إلى سلاسل الإجراءات المألوفة. على سبيل المثال، غالبًا ما يتضمن وضع الطفل في السرير لفه في الملابس وقراءة قصة وغناء أغنية له من أجل تهدئته وتشغيل الإنارة الليلية. يمكن استخدام بنى البيانات تلك في الإجابة عن الأسئلة، بل لاقتراح أسئلة أيضًا. إذا أسقطت الأم الإنارة الليلية، فيمكن إثارة أسئلة من قبيل «لماذا؟» و«ما الذي حدث بعد ذلك؟» بعبارة أخرى، هنا تكمن بذرة القصة. ومن ثم، هذا الشكل من تمثيل المعرفة يُستخدم لكتابة القصص تلقائيًا، وربما يكون ضروريًا لأجهزة الكمبيوتر «المرافقة» القادرة على الدخول في المحادثات العادية مع البشر (انظر الفصل الثالث).

الشبكات الدلالية من الأشكال البديلة لتمثيل المعرفة الخاصة بالمفاهيم (وهي عبارة عن شبكات محلية؛ انظر الفصل الرابع). في ستينيات القرن العشرين، كان روس كويليان رائدًا في هذا الشكل واتخذ نماذج للذاكرة الترابطية لدى البشر، وتتوفر الآن العديد من الأمثلة المكثفة (مثل برنامج «ووردنت») باعتبارها موارد عامة للبيانات. تربط الشبكة الدلالية المفاهيم بواسطة العلاقات الدلالية مثل «الترادف والتضاد والتبعية والشمول والعلاقة بين الجزء والكل»، كما تربطها كثيرًا عن طريق العلاقات الترابطية التي تضم المعارف الواقعية للعالم وصولًا إلى الدلالات (انظر الفصل الثالث).

يمكن للشبكة أن تمثل الكلمات وكذلك المفاهيم عن طريق إضافة روابط ترميز لكل من «المقاطع والحروف الأولية والصوتيات والمتجانسات». هذه الشبكة استخدمها نظام جيب JAPE الذي صمّمه كيم بينستيد، وبرنامج ستاند أب STAND UP الذي صمّمه جرايم ريتشي، ويولد نكات (من تسعة أنواع مختلفة) بناءً على التورية والجناس وتبديل المقاطع. على سبيل المثال: س: «ماذا تسمى القطار المكتئب؟» ج: قاطرة تعيسة؛ س: ما الذي سينتج عن التزاوج بين نعجة وكنغر؟ ج: قفاز صوفي.

تنبيه: الشبكة الدلالية ليست مماثلة للشبكات العصبية. كما سنرى في الفصل الرابع، تمثل الشبكات العصبية الموزعة المعرفة بطريقة مختلفة تمامًا. في تلك الشبكات، لا تمثل المفاهيم الفردية بعقدة مفردة في شبكة ترابطية محددة بعناية، بل باستخدام نمط يغير النشاط عبر الشبكة بالكامل. تلك الأنظمة يمكن أن تتحمل تضارب الأدلة؛ ومن ثم لا تعرقلها مشكلات الحفاظ على الاتساق المنطقي (وهذا ما سنوضحه في القسم الآتي). ولكنها لا تقدم استدلالات دقيقة. وعلى الرغم من ذلك، فهي نوع مهم للغاية في تمثيل المعرفة (وأساسٌ بالغ الأهمية للتطبيقات العملية)، وتستحق أن نُفرد لها فصلًا.

المنطق وشبكة الويب الدلالية

إذا كان الهدف النهائي لشخص ما هو الذكاء الاصطناعي العام، فيبدو أن المنطق هو أنسب اختيار ليصير تمثيلًا معرفيًا. فالمنطق يمكن تطبيقه بوجه عام. ومن حيث المبدأ، يمكن استخدام التمثيل نفسه (الرمزية المنطقية نفسها) لكل من الرؤية والتعلم واللغة وغيرها، وأي تكامل متعلق بما سبق. إضافة إلى ذلك، يوفر هذا التمثيل طرقًا قوية لإثبات النظريات من أجل التعامل مع المعلومات.

هذا يفسّر لماذا كانت طريقة التمثيل المعرفي المفضلة في بدايات الذكاء الاصطناعي هي حساب التفاضل والتكامل المُسند. وهذا الشكل من المنطق له قوة تمثيلية أكبر من منطق

القضايا؛ لأنه يمكن أن يتغلغل في الجمل للتعبير عن معانيها. على سبيل المثال، فكر في الجملة الآتية: «هذا المتجر عنده قبعة تُناسب الجميع». حساب التفاضل والتكامل المُسند يمكن أن يميز بين المعاني الثلاثة المحتملة بوضوح كما يأتي: «بالنسبة إلى كل إنسان، يوجد في المتجر قبعة تُناسبه»، و«يوجد في هذا المتجر قبعة يمكن أن يتفاوت مقاسها بحيث تناسب أي إنسان»، و«يوجد في هذا المتجر قبعة [على افتراض أنها مطوية!] بمقاس كبير لدرجة أنها تسع جميع البشر في آن واحد».

في رأي العديد من الباحثين في الذكاء الاصطناعي، لا يزال المنطق الإسنادي هو النهج المفضل. إطارات برنامج «سي واي سي» على سبيل المثال قائمة على المنطق الإسنادي. وكذلك تمثيلات معالجة اللغات الطبيعية في دلالات التراكيب (انظر الفصل الثالث). في بعض الأحيان، يمد المنطق الإسنادي أذرعه بحيث يمثل الوقت أو السبب أو الواجب/الأخلاق. بالطبع هذا يعتمد على ما إذا كان الشخص قد طوّر تلك الأشكال من منطق الموجهات أم لا؛ وهذا الأمر ليس سهلاً.

لكن المنطق له عيوب أيضاً. ومن تلك العيوب الانفجار التوافقي. كثيراً ما يستخدم المنطق طريقة «القرار»، حيث إن إثبات النظرية المنطقية قد يترسّخ في استخلاص الاستنتاجات الصحيحة، ولكنها غير مرتبطة بعضها ببعض. يوجد الاستدلال من أجل توجيه الاستنتاجات وتقييدها، ومن أجل تحديد متى ينبغي الإقلاع (وهذا ما لم يستطع صبي الساحر أن يفعله حين استحضر الجن ولم يستطع صرفه). لكن الاستدلال ليس مكفول النجاح.

عيب آخر، وهو أن إثبات نظرية القرار يفترض أن «نفي النفي إثبات». إذا كان النطاق الذي يجري التفكير فيه مفهوماً تاماً، فهذا صحيح منطقياً. لكنّ مستخدمي البرامج (مثل كثير من الأنظمة الخبيرة) المزوّدين بقرار مدمج كثيراً ما يفترضون أن عدم العثور على الضد يعني عدم وجود الضد، وهو ما يُسمّى «النفي بالفشل». وعادةً ما يكون هذا خطأً. في الحياة الواقعية، هناك فرق كبير بين إثبات أن الشيء مزيف والفشل في إثبات صحته (فكر في السؤال عمّا إذا كان شريكك يخونك أم لا).

عيب ثالث، وهو أنه في المنطق الكلاسيكي (الترتيب)، بمجرد أن تثبت صحة شيء ما، فإنه يظل صحيحاً. وعملياً، لا يبقى شيء صحيح على الدوام. قد يقبل المرء «س» لسبب وجيه (ربما كان تكليفاً افتراضياً أو حتى استنتاجاً من حجة دقيقة ودليل دامغ أو أحدهما)، ولكن قد يتبين فيما بعد أن «س» لم تُعد صحيحة، أو لم تُكن صحيحة

منذ البداية. وإذا كان الأمر كذلك، فلا بد للمرء أن يفنّد معتقداته وفقاً لذلك. بناءً على تمثيل المعرفة القائم على المنطق، فالقول أسهل من الفعل. حاول العديد من الباحثين — استلهاماً من مكارثي — أن يطوروا منطقاً «غير رتيب» يمكن أن يسع قيّم الحقيقة المتغيرة. وبالمثل، حدّد الناس العديد من أنواع المنطق «الضبابي»، حيث يمكن تسمية العبارة بأنها «محتملة/غير محتملة» أو «غير معروفة» بدلاً من «صحيحة/خاطئة». وعلى الرغم من ذلك، لم توجد آلية دفاع موثوقة ضد الرتابة.

يجد الباحثون في الذكاء الاصطناعي الذين يُطوِّرون تمثيل المعرفة القائم على المنطق في البحث عن الذرات النهائية للمعرفة — أو المعنى — بوجه عام. إنهم ليسوا الأوائل؛ مكارثي وهايز فعلاً ذلك في «بعض المسائل الفلسفية من منظور الذكاء الاصطناعي». تناولت تلك الورقة البحثية الأولى العديد من الأحجيات المعروفة، بدايةً من الإرادة الحرة وحتى الأشياء غير الواقعية. تضمّنت تلك الأحجيات أسئلة عن الكينونة الأصلية للكون؛ إذ تتضمن الحالات والأحداث والخصائص والتغيرات والإجراءات ... ماذا؟

إذا لم يكن المرء محباً لعلم ما وراء الطبيعة من صميم قلبه (شغف نادر بين البشر)، فلماذا يهتم المرء؟ ولماذا «تزيد وتيرة» السعي وراء الإجابة على تلك الأسئلة الغامضة اليوم؟ بوجه عام، الإجابة هي أن محاولة تصميم الذكاء الاصطناعي العام تثير أسئلة عن ماهية الكينونات التي يمكن أن يستخدمها تمثيل المعرفة. تثار تلك الأسئلة أيضاً في تصميم شبكة الويب الدلالية.

شبكة الويب الدلالية ليست مثل شبكة الإنترنت العالمية التي لدينا منذ تسعينيات القرن العشرين. شبكة الويب الدلالية ليست اختراعاً جديداً، بل إنها اختراع من المستقبل. إذا وُجدت أو عندما توجد، سيتحسن البحث الترابطي الموجّه بالآلة، وسيُكمّله فهم الآلة. وهذا سيمكّن التطبيقات ومتصفحات الويب من أن تصل إلى المعلومات من أي مكان على الإنترنت، وأن تدمج العناصر المختلفة بشكل معقول في التفكير بشأن الأسئلة. وهذا الأمر بالغ الصعوبة. بالإضافة إلى تطورات هندسية ضخمة في مكونات الأجهزة والبنية التحتية للاتصالات، فهذا المشروع الطموح (الذي يديره السير تيم بيرنرز-لي) يحتاج إلى تعميق فهم البرامج التي تتجول عبر الويب لما يجري.

بوجه عام، محرّكات البحث مثل جوجل وبرامج معالجة اللغات الطبيعية يمكن أن تجد ارتباطات بين الكلمات و/أو النصوص أو أحدهما، ولكن لا يوجد فهم. وهنا، هذه ليست نقطة فلسفية (انظر الفصل السادس للمزيد)، ولكنها تجريبية، كما أنها تمثّل

عقبة إضافية أمام تحقيق الذكاء الاصطناعي العام. وعلى الرغم من وجود بعض الأمثلة الخادعة المُغرية — مثل «واتسون» و«سيري» والترجمة الآلية (وكلها تناولناها في الفصل الثالث) — فإن أجهزة الكمبيوتر اليوم ليست قادرة على فهم معنى ما «يُقرأ» أو «يقال».

رؤية الكمبيوتر

بالإضافة إلى ما سبق، أجهزة الكمبيوتر اليوم لا تفهم الصور المرئية كما يفهمها البشر. (مرة أخرى، هذه نقطة تجريبية، وقد تناولنا في الفصل السادس هل أنظمة الذكاء الاصطناعي العام يمكن أن يكون لها إدراك حسي بصري واعٍ أم لا). منذ عام ١٩٨٠، اعتمدت تمثيلات المعرفة المتنوعة المستخدمة في رؤية الذكاء الاصطناعي اعتمادًا كبيرًا على علم النفس، لا سيما نظريات ديفيد مار وجيمس جيبسون. لكن على الرغم من هذه التأثيرات النفسية، فإن البرامج المرئية في الوقت الحالي محدودة للغاية.

لا أحد ينكر أن رؤية الكمبيوتر حققت إنجازات ملحوظة، ومن الأمثلة على ذلك التعرف على الوجه بنسبة نجاح تصل إلى ٩٨٪. أو قراءة خط اليد ذي الحروف المتشابكة. أو ملاحظة شخص يتصرّف بشكل يثير الريبة (باستمرار التوقف بجانب أبواب السيارات) في أماكن ركن السيارات. أو التعرف على الخلايا المريضة أفضل من اختصاصي الأمراض من البشر. عندما نصادف هذه النجاحات، فعادةً ما يُذهل العقل أيّما ذهول. لكن عادةً ما تُضطر البرامج (العديد منها عبارة عن شبكات عصبية؛ انظر الفصل الرابع) إلى معرفة ما تبحث عنه بالضبط؛ على سبيل المثال، ينبغي ألا يكون الوجه مقلوبًا، وأن يكون في صورة الصفحة الشخصية، وألا يكون مخفيًا جزئيًا خلف شيء آخر، وأن تكون الإضاءة بطريقة معيّنة (حتى تبلغ نسبة نجاح التعرف عليه ٩٨٪).

كلمة «عادةً» مهمة. في عام ٢٠١٢، أدمج مختبر الأبحاث في جوجل ١٠٠٠ جهاز كمبيوتر كبيرًا (ذا ١٦ نواة) لتكوين شبكة عصبية ضخمة تحتوي على ما يزيد على مليار نقطة اتصال. جُهزت الشبكة بتقنية التعلم العميق، وأدخل لها ١٠ ملايين صورة عشوائية من مقاطع الفيديو على موقع «يوتيوب». لم تتلقَ تعليمات بما تبحث عنه، ولم توضع علامات للصور. ولكن بعد ثلاثة أيام، تعلمت وحدة (خلية عصبية اصطناعية) أن تستجيب لصور وجوه القطط وأخرى لصور وجوه الإنسان.

هل هذا باهر؟ حسنًا، نعم. ومثير للجدل أيضًا، وسرعان ما تذكّر الباحثون فكرة «خلايا الجدة» في الدماغ. ومنذ عشرينيات القرن العشرين، اختلف علماء الأعصاب حول ما إذا كانت هذه الخلايا موجودة أم لا. والقول بوجود تلك الخلايا يعني أن هناك خلايا في الدماغ (إما خلايا عصبية مفردة وإما مجموعات خلايا عصبية صغيرة) لا تنشط أبدًا إلا عند التعرف على جدة أو سمة معينة أخرى. من الواضح أنه كان يُجري شيئًا مشابهًا في شبكة التعرف على صور القطط لدى جوجل. وعلى الرغم من أن وجوه القطط يجب أن تكون بالصورة الكاملة والوضعية الصحيحة، فإنه يمكن أن تتفاوت في الحجم أو تبدو في وضعيات مختلفة في نطاق صفيف 200×200 . دراسة أخرى دربت النظام على صور تضم صور بشر منتقاة بعناية ومختارة مسبقًا (من دون تسمية)، وبعضها من صور الملف الشخصي، ونتج عن تلك الدراسة وحدة تستطيع أن تميّز الوجوه التي حادت عن عدسة الكاميرا، ولكن لا يحدث هذا إلا في بعض الأحيان.

في الوقت الراهن، تحقّقت العديد من الإنجازات الأروع. حقّقت الشبكات المتعددة الطبقات إنجازات عظيمة في خاصية التعرف على الوجوه، ويمكنها في بعض الأحيان العثور على أبرز جزء في الصورة، وتُنشئ تعليقًا لفظيًا لوصف الصورة (مثل أشخاص يتسوّقون في متجر خارجي). انطلق في الآونة الأخيرة «تحتدي التعرف البصري الواسع النطاق»، ويزيد في كل عام عدد الفئات البصرية التي يمكن التعرف عليها، وتقل القيود المفروضة على الصور المعنية (مثل عدد الأجسام وتطابقها). ومع ذلك، ستظل أنظمة التعلم العميق ترث بعض نقاط الضعف من الأنظمة السابقة عليها.

على سبيل المثال، تلك الأنظمة شأنها شأن أداة التعرف على وجوه القطط؛ إذ لن تفهم حقيقة المساحة الثلاثية الأبعاد، وليس لديها فكرة عما يكون «الملف الشخصي» أو التطابق. حتى البرامج المصممة للروبوتات لا توفّر إلا قدرًا ضئيلاً عن تلك المسائل.

تعتمد روبوتات مارس روفر — مثل «أوبورتونينتي» و«كيوريوسيتي» — (الذين هبطا في عام ٢٠٠٤ و ٢٠١٢ على التوالي) على خدع خاصة لتمثيل المعرفة، وهي: الاستدلالات المُخصّصة للمسائل الثلاثية الأبعاد المتوقّع أن يواجهها الروبوتان. هذان الروبوتان لا يستطيعان البحث عن مسار جسم أو مناورته بوجه عام. تحاكي بعض الروبوتات الرؤية المتحركة، حيث إن حركات الجسم نفسها توفّر معلومات مفيدة (لأنها تغير المُدخلات المرئية بشكل منهجي). لكنها لا تستطيع أن تلاحظ مسارًا محتملًا، أو تتعرف على أن ذلك الشيء غير المألوف يمكن أن تلتقطه يد الروبوت في حين أنه لا يمكن ذلك.

قد يكون هناك بعض الاستثناءات بحلول وقت نشر هذا الكتاب. ولكن ستكون هناك قيود أيضًا. على سبيل المثال، لن تفهم الروبوتات عبارة «لا أستطيع التقاط هذا الشيء»؛ لأنها لن تفهم الفرق بين «أستطيع» و«لا أستطيع». يعود السبب إلى احتمالية عدم توافر منطق الموجهات لتمثيل المعرفة الخاص بها. في بعض الأحيان، قد تتجاهل الرؤية المساحة ثلاثية الأبعاد، مثل حالات قراءة خط اليد.

لكن حتى رؤية الكمبيوتر الثنائية الأبعاد محدودة. وعلى الرغم من الجهود البحثية بشأن التمثيلات «التناظرية» أو «الأيقونية»، لا يستطيع الذكاء الاصطناعي أن يستخدم المخططات بشكل موثوق في حل المسائل، مثلما نفعل في إيجاد البرهان الهندسي أو في تدوين علاقات مجردة على ظهر المغلف. (وبالمثل، لا يفهم علماء النفس حتى الآن كيف نفعل تلك الأشياء).

باختصار، معظم المآثر البصرية لدى الإنسان تفوق الذكاء الاصطناعي اليوم. وكثيرًا لا يتسم باحثو الذكاء الاصطناعي بالوضوح بشأن الأسئلة المراد طرحها. على سبيل المثال، فكر في طي رداء من الساتان الفائق النعومة طيًا مهندمًا. لا يستطيع الروبوت أن يفعل ذلك (على الرغم من أن بعضها قد يتلقّى تعليمات — خطوة خطوة — بطريقة طي منشفة منسوجة مستطيلة). أو فكر في ارتداء تيشيرت؛ يجب أن يدخل الرأس أولاً وليس الكم، ولكن لماذا؟ لا تكاد تلك المسائل الموضوعية أن يكون لها سمة بارزة في الذكاء الاصطناعي.

ولا يعني ذلك استحالة أن تبلغ رؤية الكمبيوتر مستوى رؤية الإنسان. ولكن تحقيقها أصعب مما تعتقده غالبية الناس.

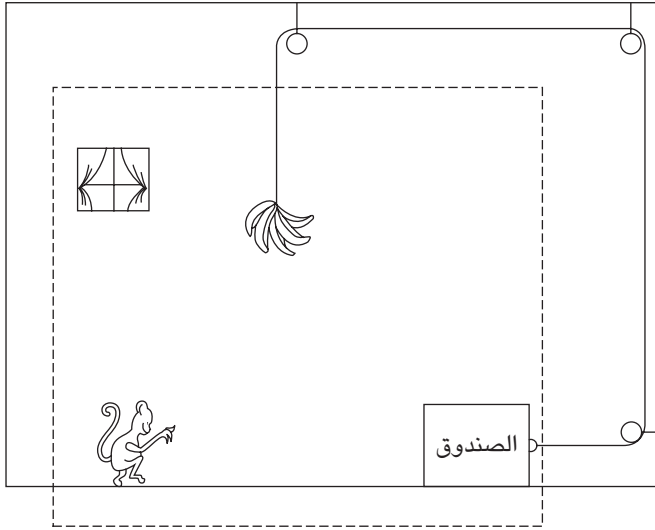
لذا، هذه حالة خاصة من الحقيقة الواردة في الفصل الأول؛ إذ تقول إن الذكاء الاصطناعي علمنا أن العقل البشري أثري بكثير وأدق مما كان يتخيله علماء النفس سابقًا. في الحقيقة، هذا هو الدرس الأساسي الذي ينبغي تعلّمه من الذكاء الاصطناعي.

مشكلة الإطار

جزء من صعوبة إيجاد تمثيل المعرفة المناسب في أي نطاق يكمن في ضرورة تجنّب «مشكلة الإطار». (انتبه: على الرغم من أن هذه المشكلة تظهر عند استخدام الأطر في صورة تمثيل المعرفة للمفاهيم، فإن مدلولات «الإطار» هنا مختلفة).

الذكاء الاصطناعي

حسب تعريف مكارثي وهايز الأساسي، تتضمن مشكلة الإطار افتراضًا (في أثناء التخطيط باستخدام الروبوت) بأن إجراء ما لن يؤدي إلا إلى «هذه» التغييرات، وفي الوقت نفسه قد يؤدي إلى «تلك» التغييرات. بوجه أعم، تظهر مشكلة الإطار عندما يتجاهل الكمبيوتر الآثار التي يفترضها الإنسان المفكر ضمنياً ولا يذكرها صراحةً. لنضرب المثال الكلاسيكي، وهو مسألة القرد والموز، وفيها يفترض القائم على حل المسألة (ربما مخطط الذكاء الاصطناعي للروبوت) أنه لا يوجد شيء ذو صلة خارج الإطار (شكل ١-٢).



شكل ١-٢: مسألة القرد والموز: كيف يصل القرد إلى الموز؟ (يفترض المنهج المعتاد لتلك المسألة أن «العالم» ذا الصلة هو الموضح خارج الإطار ذي الخط المنقط، على الرغم من عدم ذكر ذلك صراحةً. بعبارة أخرى، لا يوجد شيء خارج هذا الإطار قد يؤدي إلى تغييرات كبيرة فيه عند تحريك الصندوق).

إليكم مثالي المفضل: إذا كان الرجل بعمر ٢٠ عامًا يجمع عشرة أرطال من التوت الأسود في ساعة، والمرأة في عمر ١٨ عامًا تجمع ثمانية أرطال، فكم عدد الأرطال التي

سيجمعونها إذا تعاوننا معًا؟ بالتأكيد الإجابة ليست «١٨» رطلًا. قد تكون الكمية أكبر بكثير (لأن كليهما سيستعرض مهاراته)، أو الأرجح أنها ستكون أقل. ما أنواع المعرفة المتضمنة في هذا المثال؟ وهل يتمكن الذكاء الاصطناعي العام من التغلب على ما يبدو أنه معطيات حسابية بسيطة؟

تنشأ مشكلة الإطار لأن برامج الذكاء الاصطناعي لا تتمتع بحس «الصلة» مثل الإنسان (انظر الفصل الثالث). يمكن تجنب تلك المشكلة في حالة معرفة كل التبعات المحتملة لكل إجراء محتمل. وهكذا يكون الأمر في بعض المجالات الفنية/العلمية. ولكنه لا يكون كذلك بوجه عام. وهذا هو السبب الرئيسي في افتقار أنظمة الذكاء الاصطناعي إلى الإدراك السليم.

باختصار، تحوم مشكلة الإطار حول كل شيء، كما أنها عقبة كبيرة في طلب الذكاء الاصطناعي العام.

العوامل والإدراك الموزع

عامل الذكاء الاصطناعي إجراء قائم بذاته (مستقل)، ويُشَبَّه في بعض الأحيان بردود الفعل الانعكاسية، وفي أحيان أخرى بالعقل المصغر. يمكن أن يُطلق على تطبيقات الهواتف أو المدقق اللغوي عوامل، ولكنها ليست عوامل في العادة؛ لأن العوامل تتعاون عادةً. إنها تستخدم ذكاءها المحدود للغاية بالتعاون مع — أو على أي حال جنبًا إلى جنب مع — العوامل الأخرى كي تعطي نتائج لا يمكن أن تحققها بمفردها. التفاعل بين العوامل مهم بقدر ما هو مهم بين الأفراد أنفسهم.

تُنظَّم بعض أنظمة العوامل تنظيمًا هرميًا؛ الأفضل ثم الأدنى إذا جاز التعبير. ولكن العديد منها يجسّد الإدراك الموزع. هذا يتضمن التعاون مع بنية الأوامر الهرمية (ومن هنا كانت المراوغة — في وقت سابق — بين عبارة «بالتعاون بين» وعبارة «جنبًا إلى جنب»). لا يوجد خطة مركزية ولا تأثير من سلطة أعلى إلى سلطة أدنى، ولا يوجد أفراد يمتلكون كل المعرفة ذات الصلة.

طبيعي أن تشتمل أمثلة الإدراك الموزع الحالية على مسارات النمل وملاحة السفن والعقل البشري. تبرز مسارات النمل من سلوك العديد من فرادى النمل، حيث إنها تُخلف موادًا كيميائية لتقائماً وهي في طريقها (وتتبع تلك المواد). وبالمثل، ملاحة السفن ومناوراتها تنتج عن تضافر الأنشطة بين العديد من البشر؛ والقبطان نفسه لا يمتلك

كل المعرفة اللازمة، وبعض أفراد الطاقم لديهم قدر ضئيل من المعرفة في واقع الأمر. حتى العقل الفردي يتضمن الإدراك الموزع؛ لأنه يدمج العديد من النظم الفرعية المعرفية والتحفيزية والعاطفية (انظر الفصلين الرابع والسادس).

تتضمن الأمثلة الاصطناعية كلاً من الشبكات العصبية (انظر الفصل الرابع)، ونموذج الأنثروبولوجيا الحاسوبية الخاص بملاحة السفن وعمل الحياة الاصطناعية على الروبوتات الكائنة وذكاء السرب وروبوتات السرب (انظر الفصل الخامس)، ونماذج الذكاء الاصطناعي الرمزي للأسواق المالية (حيث تكون العوامل البنوك وصناديق التحوط والشركات المساهمة الكبرى)، ونموذج تعلّم وكيل التوزيع الذكي (LIDA) الخاص بالوعي (انظر الفصل السادس).

واضح أن الذكاء الاصطناعي العام على المستوى البشري يتضمن الإدراك الموزع.

تعلّم الآلة

يتضمن الذكاء الاصطناعي العام على المستوى البشري تعلّم الآلة أيضاً. لكن لا حاجة أن يكون التعلم شبيهاً بتعلم الإنسان. ظهر المجال من عمل علماء النفس على مفهوم التعلم والتعزيز. لكن في الوقت الراهن، يعتمد تعلّم الآلة على تقنيات رياضية مخيفة؛ لأن تمثيلات المعرفة المستخدمة تتضمن نظرية الاحتمالات والإحصاءات. (يمكن القول إن علم النفس قد تخلّف كثيراً عن الركب. بالتأكيد، تكاد الأنظمة الحديثة لتعلّم الآلة تخلو من أوجه الشبه للأفكار التي تدور في عقل الإنسان. ومع ذلك، الاستخدام المتزايد للاحتمالية بايزي في مجال الذكاء الاصطناعي يوازي النظريات الحديثة في علم النفس المعرفي وعلم الأعصاب).

مجال تعلّم الآلة في الوقت الحالي يُدرّ أرباباً غزيرة. إنه يُستخدم في جميع البيانات وفي معالجة البيانات الضخمة؛ حيث إن أجهزة الكمبيوتر العملاقة تُجري مليون مليار عملية حسابية في الثانية (انظر الفصل الثالث).

تستخدم بعض أنظمة تعلّم الآلة الشبكات العصبية. ولكن كثيراً منها يعتمد على الذكاء الاصطناعي الرمزي مدعوماً بخوارزميات إحصائية قوية. في الحقيقة، تقوم الإحصائيات بالعمل، ولا يكون للذكاء الاصطناعي التقليدي الجميل دور سوى توجيه العامل إلى مكان العمل. ومن ثمّ يعتبر بعض المحترفين تعلّم الآلة أحد علوم الكمبيوتر

والإحصاءات أو أحدهما — ولا يعتبرونه ضمن الذكاء الاصطناعي. ولكن لا يوجد فرق واضح هنا.

يتضمن تعلُّم الآلة ثلاثة أنواع شاملة، وهي: التعلم الموجَّه والتعلم غير الموجَّه والتعلم المعزز. (تأصَّلت الفروق في علم النفس، وربما تدخل آليات عصبية فسيولوجية، وينطوي التعلم المعزز — فيما بين الأنواع — على الدوبامين).

في التعلم الموجَّه، المبرمج يدرِّب النظام عن طريق تحديد مجموعة من النتائج المرجوة من مجموعة مُدخَلات (مصنَّفة بالأمثلة الدالة والأمثلة غير الدالة)، ويقدم ملاحظات عمَّا إذا تحقَّقت النتائج من عدمه. يولد نظام التعلم فرضيات عن السمات ذات الصلة. وأينما أخطأ في التصنيف، يعدِّل فرضيته طبقاً لذلك. رسائل الأخطاء المحددة البالغة الأهمية (فهي ليست مجرد تغذية راجعة عن أخطاء النظام).

في التعلم غير الموجَّه، لا يوفر المستخدم النتائج المرجوة ولا رسائل الأخطاء. يُدفع التعلم بالمبدأ القائل إن السمات المترامنة يتولد عنها توقعات بأنها ستتزامن في المستقبل. ويستخدم التعلم غير الموجَّه لاكتشاف المعرفة. ومن ثم لا يحتاج المبرمج إلى معرفة الأنماط/المجموعات الموجودة في البيانات؛ فالنظام يجدها من تلقاء نفسه.

وفي النهاية، يُدفع التعلم المعزَّز بنظائر المكافأة والعقاب؛ رسائل التغذية الراجعة تخبر النظام أن ما فعل كان جيداً أو سيئاً. وفي كثير من الأحيان، لا يكون التعزيز مجرد ازدواج، ولكن تمثله أرقام مثل الدرجات في لعبة الفيديو. قد يكون «ما فعل» قراراً فردياً (مثل حركة في لعبة) أو سلسلة قرارات (مثل حركات الشطرنج التي تبلغ ذروتها مع إماتة الملك). في بعض ألعاب الفيديو، تُحدَّث الدرجة العددية مع كل حركة. في المواقف البالغة التعقيد مثل الشطرنج، لا يشار إلى النجاح (أو الفشل) إلا بعد عدة قرارات، وأحد إجراءات تكليف الائتمان يحدّد القرارات التي من المرجَّح أن تؤدي إلى النجاح.

يفترض تعلُّم الآلة الرمزي أن تمثيل المعرفة للتعلم يجب أن يتضمن شكلاً من أشكال توزيع الاحتمالات، وصحة هذا الافتراض ليست جلية. كذلك تفترض العديد من خوارزميات التعلم أن كل متغيِّر في البيانات له القدر نفسه من توزيع الاحتمالات، وكل المتغيرات مستقلة بالتبادل؛ وهذا الافتراض عادةً ما يكون زائفاً. السبب في ذلك أن افتراض الاستقلال والتوزيع المتماثل ينطوي على العديد من النظريات الرياضية الخاصة بالاحتمالات، وتلك النظريات تعتمد عليها الخوارزميات. اعتمد علماء الرياضيات افتراض الاستقلال والتوزيع المتماثل؛ لأنه يزيد من تبسيط الرياضيات. وبالمثل، استخدام افتراض

الاستقلال والتوزيع المتماثل في الذكاء الاصطناعي يبسط مساحة البحث، وهو ما يزيد من سهولة حل المسألة.

لكن إحصاءات بايزي تتعامل مع الاحتمالات المشروطة حيث لا تكون العناصر/الأحداث مستقلة. وهنا، تعتمد الاحتمالية على أدلة التوزيع بشأن النطاق. بالإضافة إلى زيادة مستوى الواقعية في تمثيل المعرفة هذا، فإنه يتيح تغيير الاحتمالات إذا ظهر دليل جديد. ومن ثم تبرز تقنيات بايزي أكثر وأكثر في الذكاء الاصطناعي، وأيضاً في علم النفس وعلم الأعصاب. تعتمد نظريات دماغ بايزي (انظر الفصل الرابع) على الأدلة غير المستقلة وغير المتماثلة التوزيع من أجل تحفيز وصقل التعلم غير الموجه في الإدراك والتحكم في الحركة.

بناءً على العديد من نظريات الاحتمالات، يوجد العديد من الخوارزميات المناسبة لأنواع معينة من التعلم ومجموعات مختلفة من البيانات. على سبيل المثال، تقبل خوارزمية آلات المتجهات الداعمة افتراض الاستقلال والتوزيع المتماثل، ويشيع فيها استخدام التعلم الموجه، لا سيما إذا كان المستخدم يفتقر إلى معرفة مسبقة متخصصة عن المجال. تفيد خوارزميات «حقيقية الكلمات» في حال أمكن تجاهل ترتيب السمات (مثل عمليات البحث عن الكلمات وليس العبارات). وإذا أُسقط افتراض الاستقلال والتوزيع المتماثل، فإن تقنيات بايزي (آلات هيلمهولتز) يمكن أن تتعلم من أدلة التوزيع.

يستخدم معظم الاحترافيين في تعلم الآلة طرقاً إحصائية جاهزة. يحظى مبتكرو هذه الأساليب بتقدير كبير في هذا المجال؛ ففي الآونة الأخيرة، وظّف موقع «فيسبوك» منشئ خوارزمية آلات المتجهات الداعمة، وفي عام ٢٠١٣ / ٢٠١٤ وظّف موقع جوجل العديد من كبار مبتكري التعلم العميق.

التعلم العميق إنجازٌ جديد واعد وقائم على شبكات متعددة الطبقات (انظر الفصل الرابع)، وبه يُعرف على الأنماط في البيانات المدخلة على عدة مستويات هرمية. بعبارة أخرى، يكشف التعلم العميق تمثيل المعرفة المتعدد المستويات، ومنها على سبيل المثال وحدات البيكسل لوحدة كشف التباين إلى وحدة كشف الحواف إلى وحدة كشف الشكل إلى أجزاء الكائن ثم إلى الكائنات.

من الأمثلة على ذلك، وحدة كشف وجوه القطط التي برزت من بحث جوجل على «يوتيوب». مثال آخر ورد في صحيفة «نيتشر» عام ٢٠١٥، وهو خوارزمية التعلم المعزّز (خوارزمية شبكة كيو العميقة) التي استطاعت ممارسة ألعاب الأتاري الكلاسيكية ٢٦٠٠

الثنائية الأبعاد. وعلى الرغم من عدم توفير مدخلات سوى وحدات البيكسل ونقاط اللعبة (وعدم معرفة سوى عدد الإجراءات المتاحة لكل لعبة)، فقد تفوّقت على البشر بنسبة ٧٥٪ في ٢٩ لعبة من أصل ٤٩ لعبة، وتفوّقت على مختبري الألعاب الاحترافيين في ٢٢ لعبة. يبقى أن نرى إلى أي مدى يمكن توسيع هذا الإنجاز. وعلى الرغم من أن خوارزمية شبكة كيو العميقة تعثر على الاستراتيجية المثلى حيث إنها تتضمن إجراءات ذات ترتيب زمني، فإنها لا تستطيع أن تتقن الألعاب التي يتجاوز تخطيطها مدة زمنية أطول. ربما يقترح علم الأعصاب تحسينات في هذا النظام. الإصدار الحالي مستوحى من مستقبلات الرؤية هابل-ويزل؛ وهي عبارة عن خلايا في القشرة البصرية لا تستجيب إلا إلى الحركة أو لخطوط ذات اتجاه معيّن. (هذا ليس بالأمر الجلل؛ فمستقبلات هابل-ويزل ألهمت نظام باندنيمونيوم أيضاً؛ انظر الفصل الأول). الأغرب من ذلك، هذا الإصدار من خوارزمية شبكة كيو العميقة مُستوحى أيضاً من «إعادة التجربة» الذي يحدث في الحصين أثناء النوم. ومثل الحصين، يخزن نظام شبكة كيو العميقة مجموعة من النماذج أو التجارب السابقة، ويعيد تنشيطها بسرعة أثناء التعلم. هذه السمة بالغة الأهمية، ويقول المصمّمون إنه سيقع «تدهور حاد» في الأداء عند تعطيلها.

الأنظمة العامة

جزء من الإثارة التي تسبّب فيها لاعب الأتاري — إذ استحق النشر في صحيفة «نيتشر» — يكمن في السير نحو الذكاء الاصطناعي العام. خوارزمية مفردة لا تستخدم تمثيل معرفة يدويّاً، ولكنها تعلمت مجموعة كبيرة من الكفاءات في مجموعة من المهام التي تنطوي على مدخلات حسية عالية الأبعاد نسبياً. لم يسبق أن أدّى برنامجٌ ذلك. ولا حتى برنامج ألفاجو AlphaGo الذي ابتكره الفريق نفسه، والذي تغلّب على بطل العالم لي سيدول في لعبة جو Go عام ٢٠١٦. ولا حتى ألفاجو زيرو AlphaGo Zero الذي تفوّق على «ألفاجو» عام ٢٠١٧، على الرغم من أنه لم يُغدّ بأي بيانات بشأن مباريات لعبة «جو» التي لعبها البشر. تسجيلًا للموقف، في ديسمبر ٢٠١٧ أتقن «ألفا زيرو» لعبة الشطرنج أيضاً، وأتقنها بعد ممارسة اللعبة لمدة أربع ساعات ضد نفسه، وبدأ بحالات عشوائية، ولكن توفّرت له قواعد اللعبة، وقد هزم البرنامج البطل في لعبة الشطرنج «ستوك فيش»؛ إذ حقّق الفوز في ٢٨ مباراة، وتعادل في ٧٢ مباراة من أصل ١٠٠ مباراة.

لكن كما لوحظ في مستهل هذا الفصل، سيحقق الذكاء الاصطناعي العام الكامل إنجازاتٍ أكبر. وإذا كان يصعب بناء نظام مُنخَصص عالي الأداء في الذكاء الاصطناعي، فإن بناء برنامج عام في الذكاء الاصطناعي أصعب بكثير. (التعلم العميق ليس هو الحل؛ يعترف هواة التعلم العميق بالحاجة إلى نماذج جديدة لدمجها مع التفكير المنطقي المُعقّد، رمز علمي لعبارة «ليس لدينا مفتاح للحل»). وهذا يفسّر لماذا تخلّى كثير من الباحثين في الذكاء الاصطناعي عن أملهم الأول، ويلتفتون إلى مهام متنوّعة ودقيقة التحديد، وغالبًا ما كان يصحبها نجاح باهر.

من رُوّاد الذكاء الاصطناعي العام الذين حافظوا على أملهم الطموح نيويل وجون أندرسون. لقد أسّسا نظام التوجيه نحو النجاح وتحقيق الأهداف SOAR ونظام التحكم المتكيف مع التفكير والعقل ACT-R على التوالي؛ انطلق النظامان في أوائل ثمانينيات القرن العشرين وظلّ كلاهما قيد التطوير (والاستخدام) لما يقارب ثلاثة عقود بعد ذلك. لكن النظامين بالغا في تبسيط المهام؛ حيث إنهما لا يرکّزان إلا على مجموعة فرعية صغيرة من الكفاءات البشرية.

في عام ١٩٦٢، درس سيمون — زميل نيويل — مسارًا متعرجًا لنملة على أرض غير معبّدة. قال إن كل حركة عبارة عن تفاعل مباشر لموقف تدركه النملة في تلك اللحظة (تلك الفكرة الأساسية للروبوتات الكائنة؛ انظر الفصل الخامس). بعد ١٠ سنوات، ألّف نيويل وسيمون كتابًا بعنوان «حل مشكلات البشر»، وقالوا إن ذكاءنا متشابه. وطبقًا لنظريتهما النفسية، فإن الإدراك والإجراء الحركي يكتملان بالتمثيلات الداخلية (قواعد «الشرط وجواب الشرط» أو عمليات الإنتاج) المخزنة في الذاكرة، أو التي بُنيت حديثًا في أثناء حل المشكلة.

قالا: «إن الإنسان بسيط للغاية حينما يُرى على أنه نظام يتحرك». ولكن التعقيدات السلوكية الناشئة مهمة. على سبيل المثال، أظهرنا أن الأنظمة التي لا تتكون إلا من ١٤ قاعدة شرطية يمكن أن تحل المسائل الحسابية الخفية (مثل تحويل الحروف إلى أرقام من ٠ إلى ٩ في المجموع الآتي: DONALD + GERALD = ROBERT، حيث إن $D = ٥$). تتعامل بعض القواعد مع تنظيمات الهدف الأساسي/الفرعي. وبعضها يوجه الانتباه (إلى عمود أو حرف معين). وبعضها يستدعي الخطوات السابقة (النتائج المتوسطة). وبعضها يتعرف على البدايات المزيّفة. والبعض يرجع إلى البداية كي يصلح تلك البدايات.

قالا إن المسائل الحسابية الخفية تمثّل نموذجًا للهندسة الحسابية لكل السلوك الذكي؛ ومن ثَمّ كان هذا النهج مناسبًا للذكاء الاصطناعي العام. منذ عام ١٩٨٠، ابتكر

نيويل (بالتعاون مع جون ليرد وبول روزنبلوم) نظام التوجيه نحو النجاح وتحقيق الأهداف. كان الهدف من النظام أن يكون نموذجًا للمعرفة بوجه عام. يضم النظام بين ثنايا منطقته كلاً من الإدراك والانتباه والذاكرة والارتباط والاستدلال والقياس والتعلم. وقد جُمع بين الاستجابات الشبيهة باستجابات النمل (الكائنة) والتفكير مع النفس. في الحقيقة، غالباً ما ينتج التفكير عن الاستجابات المنعكسة؛ لأن التسلسل المستخدم مسبقاً للأهداف الفرعية يمكن «تجزئته» في قاعدة واحدة.

في الحقيقة، فشل نظام التوجيه نحو النجاح وتحقيق الأهداف في وضع نموذج لكل أنماط المعرفة، وجرى توسيعه فيما بعد عندما أدرك الناس بعض الفجوات. يُستخدم الإصدار الحالي في العديد من الأعراض، بدايةً من التشخيص الطبي وصولاً إلى جدولة إنتاج المصانع.

تتكون أنظمة التحكم المتكيف مع التفكير والعقلانية من أنظمة مختلطة (انظر الفصل الرابع) جرى تطويرها بالجمع بين أنظمة الإنتاج والشبكات الدالية. هذه الأنظمة تعيد تنظيم الاحتماليات الإحصائية في البيئة، كما أنها تمثل نموذجاً للذاكرة المقترنة والتعرف على الأنماط والمعاني واللغة وحل المسائل والتعلم والتخيل والتحكم الحركي الإدراكي (منذ ٢٠٠٥).

من السمات الأساسية في التحكم المتكيف مع التفكير والعقلانية التكامل مع المعرفة الإجرائية والمعرفة التوضيحية. قد يعرف شخصٌ ما أن نظرية إقليدس صحيحة من دون أن يعرف طريقة استخدامها في إيجاد البراهين الهندسية. يستطيع التحكم المتكيف مع التفكير والعقلانية أن يتعلم طريقة تطبيق حقيقة افتراضية عن طريق تكوين مئات المنتجات الجديدة التي تتحكم في استخدامها في العديد من الظروف المختلفة. إنه يتعلم أي الأهداف الرئيسية والفرعية والمتفرعة عن الأهداف الفرعية المتعلقة بعضها ببعض، وفي أي ظروف، وما النتائج التي ستتولد عن إجراء معين في ظروف متعددة. باختصار، إنه يتعلم بالتجربة. ومثل نظام التوجيه نحو النجاح وتحقيق الأهداف، فإن بإمكانه تجزئة العديد من القواعد التي غالباً ما يتم تنفيذها بطريقة متسلسلة في قاعدة واحدة. وهذا يوازي الفرق بين الطريقة التي يتبعها الخبراء والطريقة التي يتبعها المبتدئون من البشر في حل المسألة «نفسها»؛ إمّا بمنتهى السهولة وإما بشق الأنفس.

نظام التحكم المتكيف مع التفكير والعقلانية له استخدامات متنوعة. يعطي معلوم رياضيات هذا النظام ملاحظات مُخصّصة، بما في ذلك معرفة المجال ذي الصلة وبنية

الأهداف الرئيسية/الفرعية لحل المسألة. وبفضل التجزئة، تتغير أبعاد الاقتراحات بتقدم الطلاب في التعليم. تتعلق الاستخدامات الأخرى بمعالجة اللغات الطبيعية، والتفاعل بين الإنسان والكمبيوتر، والانتباه والذاكرة البشرية، والقيادة والطيران، والبحث بالصور على شبكة الويب.

تزامن نظام التوجيه نحو النجاح وتحقيق الأهداف ونظام التحكم المتكيف مع التفكير مع محاولات أخرى في مجال الذكاء الاصطناعي العام، ومنها ما يأتي: نظام «سي واي سي» الذي طوّره دوجلاس لينات. صدر نظام الذكاء الاصطناعي الرمزي هذا في عام ١٩٨٤، ولم يتوقف عن التطوير منذ ذلك الحين.

بحلول عام ٢٠١٥، تضمّن نظام «سي واي سي» ٦٢ ألف «علاقة» قادرة على ربط المفاهيم في قاعدة بياناته، بالإضافة إلى ملايين الروابط بين تلك المفاهيم. تتضمن تلك الروابط روابط دلالية وروابط حقيقية مخزّنة في شبكات دلالية كبيرة (انظر الفصل الثالث)، وعددًا لا يُحصى من حقائق الفيزياء البسيطة؛ وهي المعارف غير الصورية للظواهر الفيزيائية التي لدى جميع البشر (مثل السقوط والانسكاب). يستخدم النظام كلاً من المنطقين الرتيب وغير الرتيب، وكذلك الاحتمالات للتفكير في البيانات التي لديه. (في الوقت الحالي، تُحوّل كل المفاهيم والروابط إلى تعليمات برمجية يدويًا، ولكن يُضاف التعلم البايزي أيضًا، فهذا سيمكّن نظام «سي واي سي» من التعلم من الإنترنت).

سبق للنظام أن استخدمته العديد من الهيئات الحكومية الأمريكية، مثل وزارة الدفاع (لمراقبة الجماعات الإرهابية على سبيل المثال) والمعاهد الوطنية للصحة وبعض البنوك الكبرى وشركات التأمين. صدر الإصدار الأصغر «أوبن سايك» (OpenCyc) علنًا باعتباره مصدر خلفية للعديد من التطبيقات، كما يتوافر إصدار صغير أشمل «ريسيرش سايك» (ResearchCyc) للعاملين في مجال الذكاء الاصطناعي. وعلى الرغم من المواظبة على تحديث أوبن سايك، فإنه لا يحتوي إلا على مجموعة فرعية صغيرة من قاعدة بيانات نظام «سي واي سي»، ومجموعة فرعية صغيرة من قواعد الاستدلال. وفي النهاية، سيتوفر النظام الكامل (أو شبه الكامل) تجاريًا. ولكن قد يقع النظام في أيادٍ خبيثة إذا لم تُتخذ إجراءات معيّنة للوقاية من ذلك (انظر الفصل السابع).

وصف لينات «سي واي سي» في مجلة «إيه آي ماجازين» عام ١٩٨٦ قائلاً: «استخدام المعارف المنطقية للتغلب على الهشاشة والصعوبات في اكتساب المعرفة». وهذا يعني أن النظام كان يتعامل مع تحدي مكارثي التنبئي على وجه التحديد. واليوم، أصبح النظام

رائدًا في نمذجة التفكير «المنطقي»، وأيضًا في «استيعاب» المفاهيم التي يتعامل معها (ولا يخفى أنه حتى البرامج الرائعة في معالجة اللغات الطبيعية لا تستطيع فهمها؛ انظر الفصل الثالث).

وعلى الرغم من ذلك، ينطوي النظام على العديد من نقاط الضعف. على سبيل المثال، إنه لا يستطيع التكيف مع التعبيرات المجازية (على الرغم من أن قاعدة البيانات تحتوي على عدد كبير من التعبيرات المجازية بالطبع). إنه يتجاهل عدة أنماط من الفيزياء البسيطة. وعلى الرغم من عدم توقف معالجة اللغات الطبيعية عن التحسُّن فيه، فإنها محدودة للغاية. ومن ثم لا يتضمن رؤية. باختصار، على الرغم من أهداف نظام «سي واي سي» الموسوعية، فإنه لا يلمُّ بجميع معارف الإنسان.

إحياء الحلم

ظل نيويل وأندرسون ولينات يجِدُون في العمل طيلة ٣٠ سنة. لكن في الآونة الأخيرة، حظي الذكاء الاصطناعي العام باهتمام ملحوظ. بدأ مؤتمر سنوي في عام ٢٠٠٨، وانضمت أنظمة أخرى من المفترض أنها عامة إلى أنظمة التوجيه نحو النجاح وتحقيق الأهداف والتحكم المتكيف مع التفكير والعقلانية ونظام «سي واي سي».

في عام ٢٠١٠ على سبيل المثال، أطلق رائد تعلُّم الآلة توم ميتشل نظام متعلِّم اللغات غير المنقطع (NELL) الذي طوَّره جامعة كارنيجي ميلون. هذا النظام «ذو المنطق السليم» يبني معرفته عن طريق البحث في الويب من دون توقف (مر على ذلك سبع سنوات في وقت تأليف هذا الكتاب) وبقبول التصحيحات عبر الإنترنت من المستخدمين. إنه يستطيع عمل استنتاجات بسيطة بناءً على بياناته (غير الموسومة)؛ ومنها على سبيل المثال، الرياضي جو بلوجز يلعب كرة التنس لأنه يلعب في فريق دافيس. بدأ كيانه بعدد ٢٠٠ فئة وعلاقة (مثل السيد، ناتج عن)، ولكن بعد خمسة أعوام وسَّع كيانه وجمع ٩٠ مليون معتقد مرشح، وكل معتقد له مستوى الثقة الخاص به.

الخبر السيئ هو أن نظام متعلِّم اللغات غير المنقطع لا يعرف على سبيل المثال أنه يمكنك سحب الأجسام بسلسلة بدلاً من دفعها. في الواقع، الحس السليم المفترض لجميع أنظمة الذكاء الاصطناعي العام محدود للغاية. والادعاءات بأنه تم «حل» مشكلة الإطار السيئة مضلَّة للغاية.

أصبح لنظام متعلّم اللغات غير المنقطع نظام شقيق، وهو: متعلّم الصور غير المنقطع (NEIL). تجمع بعض برامج الذكاء الاصطناعي العام المرئية جزئياً بين تمثيل المعرفة المنطقي الرمزي والتمثيلات التناظرية أو الرسومية (وهذا التمييز وضعه آرون سلومان منذ عدة سنوات، ولكنه ليس مفهوماً فهماً جيداً حتى الآن).

بالإضافة إلى ذلك، بفضل نظام المساعد المعرفي المتعلّم والمنظم (CALO) الذي أطلقه معهد ستانفورد للأبحاث، ظهر على الساحة تطبيق «سيري» (انظر الفصل الثالث)، واشترته أبل مقابل ٢٠٠ مليون دولار أمريكي عام ٢٠٠٩. تشمل المشاريع النشطة التي يمكن مقارنتها في الوقت الحالي النظام المثير للاهتمام نموذج تعلّم وكيل التوزيع الذكي LIDA الذي طوّره ستان فرانكلين (الوارد في الفصل السادس) وبرنامج «أوبن كوج» OpenCog الذي طوّره بن جورتزل؛ إذ يتعلم البرنامج الحقائق والمفاهيم في العالم الافتراضي الثري كما يتعلم من أنظمة الذكاء الاصطناعي العام الأخرى. (نظام تعلّم وكيل التوزيع الذكي أحد نظامين عموميين يركّزان على «الوعي»؛ النظام الثاني هو التعلم الاتصالي المزوّد بتحفيز القاعدة التكيفية عبر الإنترنت «كلاريون».

في عام ٢٠١٤، انطلق مشروع أكثر حداثة في الذكاء الاصطناعي العام، ويهدف إلى تطوير «هندسة الحوسبة للكفاءات الأخلاقية في الروبوتات» (انظر الفصل السابع). بالإضافة إلى الصعوبات السالفة الذكر، سيُضطر النظام إلى مواجهة العديد من المشكلات المتعلقة بالأخلاقيات.

أي نظام على مستوى الإنسان الحقيقي لن يفعل أقل من ذلك. إذن، لا عجب من أن الذكاء الاصطناعي العام يثبت أنه صعبٌ تحقيقه.

أبعاد مفقودة

تكاد كل الأنظمة العامة اليوم تركز على «المعرفة». على سبيل المثال، يهدف أندرسون إلى تحديد «إلى أي مدى تتربط المجالات الفرعية في علم النفس المعرفي». (هل نقصد «كل» المجالات الفرعية؟ وعلى الرغم من أنه يتناول التحكم الحركي، فإنه لا يتناول اللمس أو استقبال الحس العميق، وهو ما يتجلى أحياناً في الروبوتات). سيغطي الذكاء الاصطناعي العام حقاً الدافع والعاطفة أيضاً.

أدركت قلة من علماء الذكاء الاصطناعي تلك المشكلة. كتب كلٌّ من مارفن مينسكي وسلومان كتابات ثاقبة عن معمارية الحوسبة للعقول الكاملة على الرغم من أنه لم يبنِ أيُّ منهما نموذجًا للعقل الكامل.

تناولنا نموذج سلومان «مايندر» (MINDER) الخاص بالقلق في الفصل الثالث. وعمله (بالإضافة إلى النظرية النفسية لدياترتش دورنر) ألهم برنامج «مايكروبي» MicroPsi، وهو نظام ذكاء اصطناعي عام قائم على سبعة «محفزات» مختلفة، ويستخدم السمات «العاطفية» في التخطيط واختيار الإجراءات. كذلك أثار عمله في نظام تعلُّم وكيل التوزيع الذكي الذي ذكرناه سابقًا (انظر الفصل السادس).

ولكن حتى هذه الأنظمة لا ترقى إلى الذكاء الاصطناعي العام الحقيقي. وفي بيان مينسكي لعام ١٩٥٦ بعنوان «خطوات نحو الذكاء الاصطناعي»، الذي يتنبأ بمستقبل الذكاء الاصطناعي، ذكرت العقبات وكذلك المبشرات. لكننا لم نتغلب على العديد من العقبات السابقة حتى الآن. وحسب ما سيوضح لنا الفصل الثالث، فإن الذكاء الاصطناعي العام على مستوى الإنسان لم يظهر في الأفق بعد.

الفصل الثالث

اللغة، الإبداع، العاطفة

تبدو بعض مجالات الذكاء الاصطناعي صعبة للغاية، وهي: اللغة والإبداع والعاطفة. وإن لم يتمكن الذكاء الاصطناعي من نمذجة تلك المجالات، فستذهب آمال الذكاء الاصطناعي العام أدراج الرياح.

في كل حالة، تحقق أكثر ممّا يتصوّره العديد من الناس. وعلى الرغم من ذلك، فإنه لا تزال هناك صعوبات ضخمة. ومن ثمّ لم تحدث نمذجة لهذه المجالات «البشرية» الأصلية إلا إلى حدٍّ معيّن. (ناقشنا في الفصل السادس ما إذا كانت أنظمة الذكاء الاصطناعي يمكن أن تتمتع بعاطفة أو إبداع أو فهم حقيقي أم لا. والسؤال هنا يدور حول ما إذا كانت تمتلك تلك المهارات في الظاهر أم لا).

اللغة

كثير من أنظمة الذكاء الاصطناعي يستخدم خاصية معالجة اللغات الطبيعية. معظم تلك الأنظمة يركّز على «فهم» الكمبيوتر للغة التي تُقدّم له ولا يركّز على مُخرجاته اللغوية. والسبب في ذلك أن توليد معالجة اللغات الطبيعية أصعب من قبول معالجة اللغات الطبيعية.

تتعلق الصعوبات بكلّ من المحتوى الموضوعي والشكل النحوي. على سبيل المثال، رأينا في الفصل الثاني أنه يمكن استخدام تسلسلات الإجراءات المعتادة («النصوص») باعتبارها بذرة لقصص الذكاء الاصطناعي. لكن بالنسبة إلى مسألة ما إذا كان تمثيل المعرفة في الخلفية يتضمن ما يكفي من الدوافع البشرية لتأليف قصة مثيرة، فتلك قضية أخرى. النظام المتاح تجارياً الذي يكتب ملخصات سنوية تصف الوضع المالي

المتغير لشركة ما يولد «قصصًا» مملة للغاية. الروايات التي أُنتجت باستخدام الكمبيوتر وحبكات المسلسلات التليفزيونية موجودة بالفعل، لكنها لن تفوز بأي جوائز على حسن الحبكة. (ترجمات الذكاء الاصطناعي/ملخصات النصوص التي يكتبها الإنسان قد تكون أشرى، ولكن يعود الفضل إلى المؤلفين من البشر).

بالنسبة إلى الشكل النحوي، يحتوي نثر الكمبيوتر على أخطاء نحوية في بعض الأحيان، وعادةً ما يكون غير متقن للغاية. يمكن أن يكون لسرد أنتوني ديفي المنشأ باستخدام الذكاء الاصطناعي للعبة «إكس أو» (تيك تاك تو) بنيات فقرات أساسية/فرعية تتطابق مع ديناميكيات اللعبة بطريقة مناسبة تمامًا. لكن الاحتمالات والاستراتيجيات للعبة الناعورة مفهومة تمام الفهم. ستزيد صعوبة وصف تتابع الأفكار أو أعمال الشخصيات الرئيسية بطريقة منمّقة مثل القصص التي يكتبها البشر.

بالانتقال إلى قبول الذكاء الاصطناعي للغة، بعض الأنظمة بسيطة إلى حد الملل؛ فهي لا تتطلب سوى التعرف على الكلمات الرئيسية (مثل «القوائم» في تجارة التجزئة الإلكترونية)، أو التنبؤ بالكلمات المدرجة في القواميس (مثل التكملة التلقائية التي تحدث عند كتابة رسائل نصية). هناك أنظمة أخرى أكثر تطورًا وتعقيدًا إلى حد كبير.

بعض تلك الأنظمة يتطلب التعرف على الكلام، إما على الكلمات المفردة مثل التسوق الآلي عبر الهاتف، وإما على الحديث المستمر مثل الترجمة التليفزيونية في الوقت الفعلي والتتصت على الهاتف. في الحالة الأخيرة، قد يكون الهدف انتقاء كلمات معينة (مثل كلمة «تفجير» وكلمة «جهاد») أو فهم جملة كاملة، وهذا ما يثير الاهتمام أكثر. تلك هي معالجة اللغات الطبيعية؛ وزيادة على ذلك، يجب التمييز أولاً بين الكلمات نفسها التي تتلفظ بها مختلف الأصوات وبمختلف اللهجات المحلية والأجنبية. (يتوفّر التمييز بين الكلمات مجاناً في النصوص المطبوعة). أدى التعلم العميق (انظر الفصل الرابع) إلى إحراز تقدم كبير في معالجة الكلام.

من الأمثلة الرائعة على ما يشبه فهم الجملة الكاملة الترجمة الآلية، وتجميع البيانات من المجموعات الضخمة من نصوص اللغة الطبيعية، وتلخيص المقالات في الصحف والدوريات، والإجابة عن الأسئلة ذات النطاق الحر (وكثيراً ما تُوظّف في عمليات البحث باستخدام جوجل وفي تطبيق «سيري» لأجهزة آيفون).

لكن هل يمكن لتلك الأنظمة أن تفهم اللغة؟ هل يمكن أن تتسق مع النحو على سبيل

المثال؟

في بدايات بزوغ الذكاء الاصطناعي، ظنَّ الناس أنَّ فهم اللغة يتطلب الإعراب النحوي. وبُذلت جهود كبيرة في برامج الكتابة حتى تؤدي تلك المهمة. المثال البارز هو برنامج «شردلو» SHRDLU الذي ابتكره تيرى فينوجراد في معهد ماساتشوستس للتكنولوجيا في أوائل سبعينيات القرن العشرين، وهو ما لفت انتباه عدد لا يُحصى ممَّن لم يسبق لهم السماع عن الذكاء الاصطناعي أو قالوا إن تحقيقه مستحيل.

قبل هذا البرنامج تعليمات باللغة الإنجليزية تخبر الروبوت أن يبني تراكيب مكوَّنة من كلمات متنوعة، واكتشف كيف أنه ينبغي نقل كلمات معيَّنة من أجل تحقيق الهدف. لقد ترك البرنامج أثرًا كبيرًا لعدة أسباب، وبعض تلك الأسباب ينطبق على الذكاء الاصطناعي بوجه عام. والمهم في هذا المقام هو قدرته غير المسبوقة على تعيين بنية نحوية مفصَّلة للجمل المعقَّدة، مثل الجملة الآتية: كم بيضة ستستخدمها في عمل الكعكة إذا لم تكُن تعلم أن وصفة جديتك كانت خطأ؟ (جرب!).

في الأغراض التكنولوجية، تبين أن برنامج «شردلو» محبب. تضمن البرنامج كثيرًا من الأخطاء؛ ومن ثمَّ لا يمكن لغير الباحثين ذوي المهارات العالية أن يستخدموه. حينذاك، بُنيت العديد من أدوات تحليل التراكيب النحوية، ولكن لم يكُن تعميمها على النصوص الواقعية أمرًا ممكنًا. باختصار، سرعان ما تبيَّنت الصعوبة البالغة أمام الأنظمة الجاهزة كي تحلل التراكيب المعقدة.

لم تكُن التراكيب المعقدة هي المشكلة الوحيدة. في استخدام لغة البشر، يُعد السياق والصلة بالموضوع أيضًا ذوي أهمية. ولم يكُن واضحًا ما إذا كان بإمكان الذكاء الاصطناعي معالجتهما أم لا.

أعلن تقرير اللجنة الاستشارية للمعالجة التلقائية للغات التابعة للحكومة الأمريكية لعام ١٩٦٤ أن الترجمة الآلية عملية مستحيلة. بالإضافة إلى التنبؤ بعدم رغبة كثيرين في استخدام الترجمة الآلية؛ وهو ما يجعلها غير مُجدية تجاريًا (على الرغم من أن المساعدة التي تقدمها للمترجمين البشر قد تكون مُجدية)، قال التقرير إن أجهزة الكمبيوتر ستواجه صعوبة مع التراكيب اللغوية، وستنهار أمام السياق، وفوق كل هذا ستغفل الصلة بالموضوع.

كان هذا التقرير كالصاعقة بالنسبة إلى الترجمة الآلية (حيث توقف تمويلها فعليًا بين عشية وضحاها) وإلى الذكاء الاصطناعي بوجه عام. فسَّر كثيرون التقرير على أنه يُظهر عدم جدوى الذكاء الاصطناعي. قال الكتاب الأكثر مبيعًا «أجهزة الكمبيوتر والتفكير

السليم» (الصادر عام ١٩٦١): إن الذكاء الاصطناعي كان إهدارًا للمال العام. والآن، يبدو أن أفضل خبراء الحكومات يوافقون على ذلك. وكذلك كادت جامعتان أمريكيتان أن تفتحا أقسامًا للذكاء الاصطناعي، ولكنهما ألغيا القرار.

استمر العمل في الذكاء الاصطناعي على الرغم من ذلك، وعندما أظهر برنامج «شردلو» قدرته على فهم التراكيب اللغوية بعد بضعة سنوات، بدا الأمر وكأنه إثبات منتصر للذكاء الاصطناعي التقليدي الجميل. لكن سرعان ما تسَلَّلت الشكوك؛ ومن ثَمَّ وجهت معالجة اللغات الطبيعية اهتمامها المتزايد صوب السياق بدلاً من التركيب اللغوي. أخذ بعض الباحثين السياق الدلالي على محمل الجد حتى في أوائل خمسينيات القرن العشرين. تعاملت مجموعة مارجريت ماسترمان في كامبريدج في إنجلترا مع الترجمة الآلية (واسترجاع المعلومات) باستخدام مسرد المترادفات بدلاً من القاموس. رأت المجموعة أن التركيب اللغوي «جزء سطحي للغاية وزائد في اللغة، ويغفله [مَن هم في عجلة من أمرهم]، وهم مُحَقِّقون في ذلك»، وجعلت تركيزها منصباً على مجموعات الكلمات بدلاً من الكلمات المفردة. فبدلاً من محاولة ترجمة كلمة مقابل كلمة، بحثوا في النصوص المحيطة بحثاً عن كلمات تحمل المعنى نفسه. وهذه الخاصية (عندما تعمل) تتيح ترجمة الكلمات الغامضة ترجمةً صحيحة. ومن ثَمَّ، كلمة bank يمكن ترجمتها إلى العربية على أنها «ضفة» أو «مصرف» بناءً على السياق الذي وردت فيه، مثل الحديث عن «المياه» أو عن «المال» على التوالي.

يمكن تعزيز هذا النهج السياقي القائم على مسرد المترادفات بمراعاة الكلمات التي كثيراً ما يَرِد بعضها مع بعض على الرغم من اختلاف معانيها (مثل كلمة «سك» وكلمة «ماء»). وهذا ما يحدث بمرور الوقت. وإلى جانب التمييز بين الأنواع المختلفة للتشابه المعجمي — المترادفات (فارغ/شاغر)، والأضداد (فارغ/ممتلئ)، والانتماء إلى فئة (سمكة/حيوان)، والشمول (حيوان/سمكة)، ومستوى الفئة المشتركة (سك/القد/السالمون)، والجزء/الكل (زعنفة/سمكة)، فإن الترجمة الآلية اليوم تتعرف على التوارد الموضوعي (سمكة/مياه، سمكة/ضفة، سمكة/سفن، وهكذا).

أصبح واضحاً الآن أن تناول التراكيب اللغوية المعقدة ليس ضرورياً من أجل تلخيص نصوص اللغات الطبيعية أو طرح أسئلة عنها أو ترجمتها. واليوم، تعتمد معالجة اللغة الطبيعية على قوة الإمكانيات (إمكانات الحوسبة) أكثر من اعتمادها على العقل (التحليل النحوي). تغلبت الرياضيات — لا سيما الإحصاء — على المنطق، وحلَّ تعلُّم الآلة

(على سبيل المثال لا الحصر، التعلم العميق) محل تحليل التراكيب اللغوية. هذه النهج الجديدة في معالجة اللغات الطبيعية، بدءاً من النصوص المكتوبة إلى التعرف على الكلام، تتسم بالكفاءة لدرجة أن يُتخذ معدل نجاح نسبته ٩٥٪ معياراً لقبول التطبيقات العملية. في معالجة اللغات الطبيعية الحديثة، تُجري أجهزة كمبيوتر قوية عمليات بحثٍ في مجموعات ضخمة (موسوعات) من النصوص للبحث عن أنماط الكلمات، سواء الشائعة أو غير المتوقعة (وبشأن الترجمة الآلية، هذه عبارة عن نصوص ثنائية اللغة ترجمها بشر). بإمكان تلك المعالجة أن تتعلم الاحتمالية الإحصائية لكلمات مثل «سمك/مياه»، أو «سمك/شرغوف»، أو «سمك ورقائق بطاطس/ملح وخل». (حسبما أشير في الفصل الثاني)، بإمكان معالجة اللغات الطبيعية أن تتعلم بناء «متجهات الكلمات» التي تمثل السحب المحتملة للمعنى الذي يخدم مفهوماً معيناً. لكن بوجه عام، ينصب التركيز على الكلمات والعبارات وليس التركيب اللغوي. لم يُغفل النحو؛ يمكن تعيين تسميات مثل «صفة» أو «حال» — إما تلقائياً وإما يدوياً — إلى بعض الكلمات في النصوص الخاضعة للفحص. ولكن قلما يُستخدم تحليل التراكيب اللغوية.

حتى التحليل الدلالي المفصل ليس بارزاً. يستخدم علم المعاني «الحاسوبي» التراكيب اللغوية في تحليل معاني الجمل، لكنه يوجد في المختبرات البحثية وليس في التطبيقات الواسعة النطاق. تمتلك أداة التفكير «المنطقي» «سي واي سي» تمثيلات دلالية كاملة نسبياً للمفاهيم (الكلمات) التي تحتويها؛ ومن ثم تفهمها فهماً أفضل (انظر الفصل الثاني). ولكن لا يزال هذا غير عادي.

يمكن أن تحقق الترجمة الآلية في الوقت الحالي نجاحاً باهراً. بعض الأنظمة مقيّدة بعدد صغير من الموضوعات، ولكن هناك أنظمة أخرى مفتوحة أكثر. تقدّم خدمة الترجمة من جوجل الترجمة الآلية في موضوعات غير مقيّدة لما يزيد على ٢٠٠ مليون مستخدم كل يوم. نظام «سيستران» يستخدمه الاتحاد الأوروبي (لأكثر من ٢٤ لغة) وحلف الناتو وشركة زيروكس وجنرال موتورز.

توشك العديد من هذه الترجمات، بما فيها وثائق الاتحاد الأوروبي، أن تبلغ حد الكمال (لأنه لا تُستخدم سوى مجموعة فرعية من الكلمات في النصوص الأصلية). لكن كثيراً من الترجمات غير مثالي، ولكن يسهل فهمه لأن القارئ المطلع يمكنه تجاهل الأخطاء النحوية واختيارات الكلمات غير الموفقة — مثلما يفعل المرء عندما يستمع إلى متحدث لا يتحدث بلغته الأصلية. وبعض النصوص تتطلب أن يُجري المترجمون تعديلات طفيفة بعد

الترجمة. (في اللغة اليابانية، قد يتطلب الأمر تعديلاً كبيراً قبل الترجمة وبعدها. لا تحتوي اللغة اليابانية على كلمات مجزأة، مثل الفعل vot-ed في الزمن الماضي في اللغة الإنجليزية حيث يضاف المقطع ed، كما يُعكس ترتيب العبارات. وعادةً ما يصعب المطابقة الآلية للغات من مجموعات اللغات المختلفة).

باختصار، عادةً ما تكون نتائج الترجمة الآلية جيدة بالدرجة التي يفهمها المستخدم البشري. وبالمثل، برامج معالجة اللغات الطبيعية «أحادية اللغة» التي تلخص المقالات الصحفية يمكن أن تُبين هل المقال يستحق القراءة بالكامل أم لا. (يمكن القول إن الترجمة المثالية مستحيلة على أي حال. على سبيل المثال، يتطلب طلب تفاحة في اليابانية لغةً تُبرز الوضع الاجتماعي المقارن للمحاورين، ولكن لا توجد فروق مكافئة في اللغة الإنجليزية). أما الترجمة الآلية في الزمن الحقيقي، والتي تُتاح على بعض تطبيقات الذكاء الاصطناعي مثل «سكايب»، فهي أقل نجاحاً. يرجع السبب في ذلك إلى أن النظام ينبغي أن يتعرف على الكلام وليس على نصوص مكتوبة (حيث إنه يُفصل بين الكلمات الفردية فصلاً واضحاً).

ظهر تطبيقان بارزان آخران من تطبيقات معالجة اللغات الطبيعية بأشكال استرجاع المعلومات، وهما: «البحث الموزون» (الذي أطلقته مجموعة ماسترمان عام ١٩٧٦) و«التنقيب عن البيانات». على سبيل المثال، محرك البحث جوجل يبحث عن المصطلحات الموزونة بمدى الصلة بالموضوع — ويتم التقييم إحصائياً وليس دلاليّاً (أي من دون فهم). يبحث التنقيب عن البيانات عن أنماط الكلمات التي لا يشك فيها المستخدم البشري. يُستخدم التنقيب عن البيانات منذ مدة طويلة في عمليات البحث عن المنتجات والعلامات التجارية في الأسواق، ويطبّق الآن (غالباً باستخدام التعلم العميق) على «البيانات الضخمة»: وتعني عمليات الجمع الضخمة للنصوص (متعددة اللغات في بعض الأحيان) أو الصور، مثل التقارير العلمية أو السجلات الطبية أو المدخلات على مواقع التواصل الاجتماعي والإنترنت.

تتضمن تطبيقات التنقيب عن البيانات الضخمة المراقبة ومكافحة الجاسوسية ورصد المواقف العامة للحكومات وصناع السياسات وعلماء الاجتماع. تقارن هذه الاستفسارات بين الآراء المتغيرة للمجموعات الفرعية المتميزة: الرجال/النساء، الشباب/الكبار، الشمال/الجنوب، وهكذا. على سبيل المثال، حلّ مركز الأبحاث البريطاني ديموس بالتعاون مع فريق تحليل بيانات معالجة اللغات الطبيعية بجامعة ساسكس) آلاف

الرسائل على موقع «تويتر» المتعلقة بكراهية النساء والمجموعات العرقية والشرطة. يمكن البحث عن التدفقات المفاجئة من التغريدات بعد أحداث معينة («حوادث تويتر») لاكتشاف التغييرات في الرأي العام حول رد فعل الشرطة على حادثة معينة على سبيل المثال.

يبقى أن نرى هل معالجة اللغات الطبيعية في البيانات الضخمة سيؤدي بالضرورة إلى نتائج مفيدة أم لا. وفي كثير من الأحيان، لا يسعى التنقيب عن البيانات (باستخدام التحليل الدلالي) إلى قياس مستوى الاهتمام العام فحسب، بل أيضًا إلى قياس النبرة التقييمية للعامة. ولكن الكلام لا يكون صريحًا. على سبيل المثال، قد تحتوي تغريدة على كلام عنصري مهين في الظاهر؛ ومن ثم تشقُّره الآلة على أنه «كلام سلبي» من حيث الدلالة، وهو ليس مهينًا في حقيقته. وقد يحكم إنسان حينما يقرأ التغريدة أن المصطلح المستخدم (في هذه الحالة) علامة إيجابية بالنسبة إلى هوية جماعة، أو أنه وصف محايد (مثل متجر «باكي» عند الزاوية) وليس إهانة أو إساءة. (توصَّل البحث التجريبي إلى أنه لا توجد سوى نسبة ضئيلة من التغريدات التي تنطوي على مصطلحات عرقية/إثنية عدائية بالفعل).

وفي تلك الحالات، سيعتمد حكم الإنسان على السياق، كأن يعتمد على الكلمات الأخرى في التغريدة. قد يكون من الممكن تعديل معايير البحث الخاصة بالآلة بحيث يقل الانتساب إلى «المشاعر السلبية». ثم مرة أخرى، قد لا يكون الأمر كذلك. فكثيرًا ما تكون هذه الأحكام مثيرة للخلاف. حتى في وقت الاتفاق، قد يصعب تحديد أنماط السياق التي تبرر تفسير الإنسان.

هذا مجرد مثال واحد على صعوبة تحديد مدى الصلة بالموضوع من حيث المنظور الحاسوبي (أو حتى الشفهي).

هناك تطبيقان معروفان في معالجة اللغات الطبيعية قد يبدوان للوهلة الأولى أنهما يتعارضان مع هذا الكلام، وهما: تطبيق «سيري» من أبل، وتطبيق «واتسون» من آي بي إم.

تطبيق «سيري» هو مساعد شخصي (قائم على القواعد)، و«روبوت دردشة» متحدث يمكنه الإجابة عن العديد من الأسئلة المختلفة بسرعة. التطبيق يمكنه الوصول إلى كل شيء على الإنترنت — بما في ذلك خرائط جوجل وويكيبيديا وجريدة «نيويورك تايمز» التي لا يتوقَّف تحديثها وقوائم الخدمات المحلية مثل الضرائب والمطاعم. التطبيق يمكنه الوصول أيضًا إلى الموقع القوي للإجابة عن الأسئلة «وولفارم ألفا» الذي يمكنه استخدام التفكير

المنطقي لاستنباط الإجابات عن مجموعة كبيرة من الأسئلة الحقيقية، وليس مجرد البحث عن تلك الإجابات.

يقبل تطبيق «سيري» الأسئلة المنطوقة من المستخدم (حيث إنه يتكيف تدريجياً مع صوت المستخدم ولهجته)، ويجب عنها باستخدام البحث على شبكة الإنترنت وتحليل المحادثات. يدرس تحليل المحادثات كيف يرتّب الناس تسلسل الموضوعات في المحادثة، وكيف ينظّمون التفاعلات مثل الشرح والاتفاق. هذا النهج يُمكن تطبيق «سيري» من التفكير في الأسئلة مثل «ما الذي يريده المحاور؟» و«ما الإجابة التي ينبغي أن يجيب بها؟» حتى إنه يتكيف مع اهتمامات المستخدم الفردية وتفضيلاته.

باختصار، يبدو أن تطبيق «سيري» ليس حساساً تجاه الصلة بالموضوع فحسب، بل تجاه الصلة بالتفضيلات الشخصية أيضاً. ومن ثمّ التطبيق مثير للإعجاب في ظاهره. ومع ذلك، يسهل أن يعطي التطبيق إجابات تافهة — وإذا شرد المستخدم عن نطاق الحقائق، يضل تطبيق «سيري» الإجابة.

وتطبيق «واتسون» أيضاً يركّز على الحقائق. وباعتباره مورداً جاهزاً (يحتوي على ٢٨٨٠ معالجا أساسياً) لمعالجة البيانات الضخمة، والتطبيق مستخدم بالفعل في بعض مراكز الاتصال، ويجري تكييفه مع بعض الاستخدامات الطبية مثل تقييم علاجات السرطان. ولكنه لا يجيب عن الأسئلة المباشرة فحسب مثل تطبيق «سيري». يمكنه أيضاً التعامل مع الألغاز التي تنشأ في لعبة المعرفة العامة «جيوباردي» (Jeopardy).

في لعبة «جيوباردي»! لا تُطرح أسئلة مباشرة على اللاعبين، بل تعطى لهم تلميحات وعليهم التخمين عمّ يكون السؤال المرتبط بذلك التلميح. على سبيل المثال، يُقال لهم «في ٩ مايو ١٩٢١، هذه الخطوط الجوية «المعروفة عن ظهر قلب»، فتحت أول مكتب للسفر في أمستردام»، وعلى اللاعب أن يجيب عن السؤال الآتي: «ما الخطوط الجوية الملكية الهولندية؟»

باستطاعة تطبيق «واتسون» أن يجيب عن هذا السؤال والعديد من الأسئلة الأخرى. وعلى خلاف تطبيق «سيري»، فإن إصدار اللعبة «جيوباردي»! ليس له اتصال بالإنترنت (على الرغم من أن الإصدار الطبي له اتصال بالإنترنت)، ولم يُغذَّ بأي فكرة عن بنية المحادثات. ولا يمكنه اكتشاف الإجابة عن طريق التفكير المنطقي. بل إنه يستخدم البحث الإحصائي المتوازي على نطاق واسع عبر قاعدة بيانات هائلة ولكنها مغلقة. تحتوي تلك القاعدة على مستندات — مراجعات وكتب مرجعية لا حصر لها، بالإضافة إلى مجلة «نيويورك تايمز»، حيث إنها تقدم حقائق عن الجذام إلى لست وحقائق عن الهيدروجين

إلى هيدرا، وهكذا. عند ممارسة لعبة «جيوباردي!» تُوجَّه عمليات البحث فيه بمئات الخوارزميات ذات التصميم الخاص، والتي تُبرز الاحتمالات الكامنة في اللعبة. ويمكن أن تتعلم من تخمينات المتسابقين البشريين.

في عام ٢٠١١، نافس تطبيق «واتسون» كمبيوتر «ديب بلو» من شركة آي بي إم الذي نافس كاسبروف (انظر الفصل الثاني)؛ إذ واضح أنه هزم أفضل بطلين بشريين. («واضح» لأن الكمبيوتر يتفاعل على الفور في حين أن الإنسان يحتاج بعض الوقت للتفاعل قبل أن يضغط على الجرس). ولكن كما حدث مع كمبيوتر ديب بلو، لم يحز الفوز في كل مرة. في إحدى المرات، خسر لأنه على الرغم من تركيزه بشكل صحيح على ساق رياضية معينة، فإنه لم يدرك أن الحقيقة الحاسمة في بياناته المخزنة هي أن هذا الشخص قد فقد ساقه. لن يتكرر هذا الخطأ لأن مبرمجي «واتسون» أشاروا إلى أهمية كلمة «مفقود». لكن ستحدث أخطاء أخرى. حتى في سياقات البحث عن الحقائق العادية، غالبًا ما يعتمد الأشخاص على أحكام ذات صلة تفوق قدرات «واتسون». على سبيل المثال، تطلبت إحدى الإشارات هوية اثنين من حواربي المسيح، واسمهما من أشهر أسماء الأطفال، وينتهيان بالحرف نفسه. كانت الإجابة «ماثيو وأندريو»، وقد توصل تطبيق «واتسون» إلى الإجابة فورًا. كما توصل البطل البشري إلى الإجابة هو الآخر. ولكن راوده في البداية الاسمان «جيمس وجوداس». نبذ تلك الإجابة، وتذكر أن السبب الوحيد هو «أنه لا يعتقد أن اسم جوداس مشهور بين الأطفال لسبب أو لآخر». لم يفعل تطبيق «واتسون» ذلك.

غالبًا ما تكون الأحكام البشرية ذات الصلة بالموضوع أقل وضوحًا من ذلك بكثير، وتكون دقيقة للغاية بالنسبة إلى معالجة اللغات الطبيعية الحالية. في الحقيقة، الصلة بالموضوع نسخة لغوية/مفاهيمية من «مشكلة الإطار» المستعصية في الروبوتات (انظر الفصل الثاني). ويقول كثيرون إنه لن يتمكّن نظام غير بشري من إتقانها إتقانًا تامًا. سنناقش في الفصل السادس ما إذا كان ذلك نابغًا من التعقيد الهائل وحده أم من حقيقة أن ربط الموضوعات متأصل في تكويننا البشري على وجه التحديد.

الإبداع

الإبداع هو القدرة على إنتاج أفكار أو أعمال فنية جديدة ومثيرة للدهشة وقيمة، وهو قمة الذكاء البشري، كما أنه ضروري من أجل الذكاء الاصطناعي العام على مستوى الإنسان. لكن يراه كثيرون على أنه شيء غامض. ليس واضحًا كيف تبادر الأفكار الجديدة إلى عقل الإنسان، ناهيك عن أجهزة الكمبيوتر.

حتى إدراك تلك الأفكار ليس مباشرًا؛ فغالبًا ما يختلف الناس هل الفكرة مبدعة أم لا. تدور بعض الخلافات حول ما إذا كانت الفكرة جديدة بالفعل أم لا وبأي معنى. قد لا تكون الفكرة جديدة إلا على الفرد المعني، أو حتى جديدة على تاريخ البشرية بالكامل (بحيث تجسد الإبداع «الفردية» و«التاريخية» على التوالي). وعلى أي حال، قد يزيد وجه التشابه أو يقل مع أفكار سابقة، وهو ما يترك مساحة لمزيد من الاختلافات). تدور الخلافات الأخرى حول القيمة (التي تتضمن الوعي الوظيفي والاستثنائي في بعض الأحيان؛ انظر الفصل السادس). يمكن للفكرة أن تقيّمها فئة اجتماعية دون غيرها. (لنضرب مثالًا بالازدراء الذي يوجهه الشباب اليوم لأي شخص يعتز بأقراص «دي في دي» من أبا).

يفترض عمومًا أن الذكاء الاصطناعي ليس لديه أي شيء مثير للاهتمام يقوله عن الإبداع. لكن تكنولوجيا الذكاء الاصطناعي ولدت كثيرًا من الأفكار الجديدة تاريخيًا والمثيرة للدهشة والقيمة. تنشأ هذه الأفكار على سبيل المثال في تصميم المحركات والمستحضرات الصيدلانية وأنواع مختلفة من فنون الكمبيوتر.

إضافةً إلى ذلك، تساعد مفاهيم الذكاء الاصطناعي في شرح الإبداع البشري. فهي تمكّننا من التمييز بين ثلاثة أنواع، وهي: الإبداع التوافقي والاستكشافي والتحويلي. تتضمن تلك الأنواع آليات نفسية مختلفة، وهو ما يثير أنواعًا من المفاجآت.

في الإبداع التوافقي، يجري توفيق الأفكار المعتادة بطرق غير معتادة. تتضمن الأمثلة الملصقات المرئية والصور الشعرية والتشبيهات العلمية (تشبيه القلب بالمضخة والذرة بالمجموعة الشمسية). يخلق التوفيق الجديد مفاجأة إحصائية؛ فقد كان الأمر مستبعدًا كلاعب فرصته ضعيفة في مسابقة ولكنه فاز فيها. لكن التوفيق مفهوم وقيم للغاية. يعتمد «مستوى القيمة» على الأحكام ذات الصلة التي ناقشناها مسبقًا.

الإبداع الاستكشافي تقل فيه سمة التميز؛ حيث إنه يستخدم بعض طرق التفكير ذات التقدير الثقافي (مثل أساليب الرسم أو الموسيقى، أو المجالات الفرعية للكيمياء أو الرياضيات). تُستخدم القواعد الأسلوبية (من دون وعي إلى حد كبير) للخروج بالفكرة الجديدة، مثلما يتولد عن قواعد النحو جمل جديدة. قد يستكشف الرسام/العالم إمكانات الأسلوب بطريقة لا تترك مجالًا للتفنيد. أو قد يعتمدون التشجيع عليه واختباره، بحيث يكتشفون ما يمكن أن ينتج عنه وما لا يمكن أن ينتج عنه. يمكن أيضًا إدخال تعديلات على الأسلوب عن طريق التغيير الطفيف في قاعدة ما (مثل الإضعاف/التقوية). وعلى الرغم

من حادثة التركيب الجديد، فإنه سينظر إليه على أنه يندرج ضمن مجموعة الأساليب المألوفة.

الإبداع التحويلي خليفة الإبداع الاستكشافي، وعادةً ما ينشأ عن التخلص من قيود الأسلوب الحالي. وهنا يتغير قيد أو أكثر من قيود الأساليب تغيراً جذرياً (بالحذف، النفي، التكملة، التعويض، الإضافة ...) ومن ثم تنشأ تراكيب جديدة لم يكن إنشاؤها ممكناً من قبل. تلك الأفكار الجديدة تثير الدهشة في عمقها؛ لأنها في ظاهرها تبدو «مستحيلة». وغالباً ما تكون غير مفهومة في البداية، حيث إنها لا تُفهم فهماً تاماً من حيث طريقة التفكير التي كانت مقبولة من قبل. ولكن يجب أن يكون قريباً من طريقة التفكير السابقة واضحاً إذا كان يُبتغى قبولها. (في بعض الأحيان، يستغرق هذا الاعتراف عدة سنوات). كل أنواع الإبداع الثلاثة تحدث في الذكاء الاصطناعي — ولكن كثيراً ما ينسب الراصدون النتائج إلى الإنسان (باجتياز اختبار تورينج فعلياً؛ انظر الفصل السادس). ولكنها ليست بالنسب التي قد يتوقعها المرء.

وعلى وجه الخصوص، عدد الأنظمة التوافقية ضئيل للغاية. قد يعتقد المرء أنه سهل نمذجة الإبداع التوافقي. وفي النهاية، لا شيء أبسط من جعل الكمبيوتر ينتج ارتباطات غير مألوفة من الأفكار المخزنة بالفعل. غالباً ما ستكون النتائج جديدة تاريخياً، ومدهشة (إحصائياً). ولكن إذا كان يُبتغى تقديرها، فلا بد أن تكون ذات صلة بعضها ببعض. وتلك السمة ليست صريحة كما رأينا. فبرامج تأليف النكات التي ذكرناها في الفصل الثاني تستخدم قوالب نكات كي تساعد في تأليف نكات ذات صلة. وبالمثل، يبني التفكير المنطقي القائم على الحالات في الذكاء الاصطناعي الرمزي نظائر بفضل الأمثلة المماثلة التركيبية المحولة إلى تعليمات برمجية مسبقاً. ومن ثم يتضمن الإبداع التوافقي فيها مزيجاً قوياً من الإبداع الاستكشافي كذلك.

وعلى النقيض من ذلك، قد يتوقع المرء أن الذكاء الاصطناعي لا يستطيع أبداً أن يصوغ نموذجاً للإبداع التحويلي. هذا التوقع أيضاً مغلوط. بالتأكيد لا يستطيع البرنامج فعل شيء أعلى من قدراته. لكن البرامج التطورية يمكنها أن تعدّل من نفسها (انظر الفصل الخامس). يمكنها أيضاً تقييم أفكارها التي عدّلتها حديثاً، لكن هذا لا يحدث إلا إذا وفّر المبرمج معايير انتقاء واضحة. عادةً ما تُستخدم هذه البرامج في تطبيقات الذكاء الاصطناعي التي تسعى إلى التجديد، مثل تصميم أدوات علمية أو أدوية جديدة.

لكن هذا الطريق ليس طريقاً ساحراً إلى الذكاء الاصطناعي العام. فنادراً ما تتوافر ضمانات لنتائج قيمة. يمكن الاعتماد في بعض البرامج التطورية (في الرياضيات أو العلوم)

في التوصل إلى الحل الأمثل، ولكن لا يمكن شرح العديد من المسائل بإدخال تحسينات في البرنامج. الإبداع التحويلي محفوف بالمخاطر بسبب كسر القواعد التي سبق قبولها. يجب تقييم أي تراكيب جديدة، وإلا عمّت الفوضى. ولكن وظائف اللياقة في الذكاء الاصطناعي اليوم يحددها البشر؛ فالبرامج لا تستطيع أن تعدلها/تطورها من تلقاء نفسها.

الإبداع الاستكشافي أفضل نوع يتناسب مع الذكاء الاصطناعي. الأمثلة لا حصر لها. مُنحت براءات اختراع لبعض أفكار الذكاء الاصطناعي الاستكشافي الجديدة في الهندسة (بما في ذلك فكرة تولدت عن برنامج طوره مصمم مشروع «سي واي سي»؛ انظر الفصل الثاني). وعلى الرغم من أن الفكرة الحاصلة على براءة اختراع «ليست واضحة لشخص ماهر في الفن»، فإنها قد تكمن بشكل غير متوقع في إمكانات الأسلوب الذي يجري استكشافه. لا يمكن تمييز بعض استكشافات الذكاء الاصطناعي عن إنجازات الإنسان الرائعة، مثل الألحان الموسيقية التي تؤلفها برامج ديفيد كوب بأسلوب الموسيقى شوبان أو الموسيقى باخ. (كم مرة يستطيع الإنسان فعل ذلك؟)

لكن حتى الذكاء الاصطناعي الاستكشافي يعتمد اعتمادًا أساسيًا على الحكم البشري. فالشخص يجب أن يميز القواعد الأسلوبية ويذكرها بوضوح. وعادةً ما يكون هذا صعبًا. أحد الخبراء العالميين في منازل البراري التي صممتها مؤسسة فرانك لويد رايت تخلى عن محاولته لوصف الأسلوب المعماري لتلك المنازل، مصرحًا بأنه «يكتنفه الغموض». وفي وقت لاحق، ولدت «قواعد الشكل» القابلة للحوسبة عددًا غير محدود من تصميمات منازل البراري، ومن ضمنها التصميمات الأصلية التي يزيد عددها على ٤٠ تصميمًا؛ ولا توجد تصميمات غير معقولة. ولكن في النهاية كان المحلل البشري هو المسئول عن نجاح النظام. وإذا تمكّن الذكاء الاصطناعي العام من تحليل الأساليب (في الفنون والعلوم) من تلقاء نفسه، فستكون الاستكشافات الإبداعية بمثابة عمله الخالص كليةً. وعلى الرغم من الأمثلة الحديثة — المحدودة للغاية — على الأساليب الفنية التي تعرف عليها التعلم العميق (انظر الفصلين الثاني والرابع)، فإن هذا الأمر بالغ الصعوبة.

مكّن الذكاء الاصطناعي الفنانين البشر من تطوير شكل جديد من أشكال الفن يُسمى الفنون المصممة باستخدام الكمبيوتر. تشمل تلك الفنون كلاً من العمارة والجرافيك والموسيقى وتصميم الرقصات بالإضافة إلى الأدب، غير أن نسبة نجاحه أقل (نظرًا للصعوبات التي تواجه معالجة اللغات الطبيعية بشأن التراكيب والصلة بالموضوع). في الفنون المصممة باستخدام الكمبيوتر، ومقارنة بفرشاة رسم جديدة، فإن جهاز الكمبيوتر

ليس مجرد أداة تساعد الفنان على إنجاز أعمال قد ينجزها على أي حال. بل إن العمل قد لا يُنجز أو ربما لا يُتخيل من دون الكمبيوتر.

تتضمن الفنون المصممة باستخدام الكمبيوتر أمثلة من أنواع الإبداع الثلاثة. ولأسباب شُرحَت من قبل، نادرًا ما يندرج أي فن من الفنون المصممة باستخدام الكمبيوتر ضمن الإبداع التوافقي. (أداة «بينتينج فول»، التي ابتكرها سيمون كولتون، صممت فنونًا تصويرية بصرية لها علاقة بالحرب، ولكن أُعطيت تعليمات مخصصة للأداة كي تبحث عن الصور المرتبطة بـ «الحرب»، وكانت الصور متاحة بالفعل وجاهزة في قاعدة بياناتها). ومن ثَمَّ يندرج معظم تلك الفنون ضمن الإبداع الاستكشافي أو التحويلي.

في بعض الأحيان، يصمّم الكمبيوتر العمل الفني بأكمله تصميمًا مستقلًا عن طريق تنفيذ التعليمات البرمجية التي يكتبها الفنان. ومن ثَمَّ تُنتج أداة آرون (AARON)، التي ابتكرها هارولد كوهين، رسوماتٍ خطية وصورًا ملوّنة من دون مساعدة (وفي بعض الأحيان تصمم ألوانًا باهرة الجمال لدرجة أن يقول كوهين إنها أفضل من تلوينه هو نفسه).

بالمقارنة بالفنون التفاعلية، فإن الشكل النهائي للعمل الفني يعتمد جزئيًا على المدخلات من الجمهور ممّن قد يكون لهم أو قد لا يكون لهم تحكّم متعمد على ما يحدث. يرى بعض الفنانين التفاعليين الجمهور كأنهم شركاء في العمل الفني، ويраهم البعض الآخر على أنهم عوامل سببية يؤثرون من دون علم في العمل الفني بعدة طرق (والبعض يراهم بالمنهجين كليهما مثل إرنست إدموندس). أما في الفنون التطورية التي من أمثلتها ويليام لاثام وجون ماكورماك، فدائمًا ما تُصمم/تُعدل النتائج باستخدام الكمبيوتر، ولكن عادةً ما يكون «الاختيار» من الفنان أو الجمهور.

باختصار، إبداع الذكاء الاصطناعي له العديد من التطبيقات. وفي بعض الأحيان، يمكن أن تضاهي أو حتى تتفوق على معايير الإنسان في زاوية صغيرة من زوايا العلوم والفنون. ولكن مضاهاة إبداع الإنسان «في العموم» مسألة مختلفة تمامًا. أصبح الذكاء الاصطناعي العام بعيدًا أكثر من أي وقت مضى.

الذكاء الاصطناعي والعاطفة

كما هي الحال مع الإبداع، عادةً ما يُنظر إلى العاطفة على أنها غريبة تمامًا عن الذكاء الاصطناعي. وإلى جانب اللامعقولية الحدسية، يبدو أن حقيقة اعتماد الأمزجة والعواطف

على المعدلات العصبية المنتشرة في الدماغ تستبعد نماذج الذكاء الاصطناعي الخاصة بالتأثير.

على مدار عدة سنوات، بدأ أن علماء الذكاء الاصطناعي أنفسهم يتفوقون على ذلك. تجاهل العلماء العاطفة، ولكن هناك بعض الاستثناءات المبكرة في ستينيات وسبعينيات القرن العشرين، وبالتحديد هربرت سيمون؛ إذ رأى أن العاطفة داخلة في التحكم المعرفي، وكينيث كولبي الذي صمّم نماذج مثيرة للاهتمام تعبر عن الاضطراب العصبي وجنون العظمة، ولكنها مُفرطة الطموح.

اليوم، أصبحت الأمور مختلفة. جرت محاكاة التعديل العصبي (في شبكات «جاس نت»؛ انظر الفصل الرابع). والآن، تعالج العديد من مجموعات البحث في الذكاء الاصطناعي مسألة العاطفة. كذلك معظم تلك البحوث (وليس جميعها) ضحلة من الناحية النظرية. ومن المحتمل أن يكون معظمها مربحاً؛ حيث إنها تهدف إلى تطوير «الرفاق الحاسوبيين». صُممت أنظمة الذكاء الاصطناعي هذه بحيث تتفاعل مع الناس بطرق مريحة عاطفياً (إلى جانب كونها مفيدة عملياً) ومُرضية للمستخدم، وبعض تلك الأنظمة قائم على الشاشة، وبعضها عبارة عن روبوتات متنقلة. معظم تلك الأنظمة مُوجّه لكبار السن والمعاقين، بمن فيهم المصابون بالخرف الأوّل. وبعضها يستهدف الأطفال والرُضع. وبعضها عبارة عن «ألعاب تفاعلية للكبار». باختصار، أدوات تقديم الرعاية الحاسوبية والمربيات الروبوتية وشركاء الجنس.

تشمل التفاعلات بين الإنسان والحاسوب المعنية ما يأتي: تقديم تذكيرات بشأن التسوق والأدوية والزيارات العائلية، والتحدث عن اليوميات الشخصية المستمرة والمساعدة في تجميعها، وجدولة البرامج التلفزيونية ومناقشتها بما في ذلك الأخبار اليومية، وصنع/إحضار الطعام والشراب، ومراقبة العلامات الحيوية (وبكاء الأطفال)، والتحدث والتحرك بطرق مثيرة جنسياً.

تتضمن العديد من تلك المهام العاطفة من جانب الإنسان. بالنسبة إلى الرفيق ذي الذكاء الاصطناعي، فقد يتمكّن من إدراك العواطف لدى المستخدم الإنسان، أو ربما يتفاعل بطرق عاطفية بشكل واضح. على سبيل المثال، قد يؤدي الحزن من جانب المستخدم — ربما بسبب ذكر فاجعة — إلى إظهار بعض التعاطف من جانب الجهاز. بإمكان أنظمة الذكاء الاصطناعي أن تتعرف على عواطف الإنسان بعدة طرق مختلفة. بعض تلك الطرق فسيولوجية؛ بمعنى أنها ترصد معدل تنفس الشخص واستجابة الجلد

الكهربية. البعض الآخر لفظي؛ بمعنى أنها تلاحظ سرعة المتحدّث ونبرة صوته، وكذلك المفردات التي يتفوّه بها. وهناك طرق بصرية؛ بمعنى أن تحلّل تعبيرات الوجه. في الوقت الراهن، كل تلك الطرق بسيطة نسبيًا. فعواطف المستخدم يسهل فقدانها، وكذلك يسهل تفسيرها تفسيرًا خاطئًا.

عادةً ما يكون الأداء العاطفي من جانب الرفيق الحاسوبي أداءً لفظيًا. إنه يعتمد على المفردات (ونبرة الصوت إذا كان النظام يولّد كلامًا). ولكن بقدر ما يراقب النظام الكلمات الرئيسية المألوفة من المستخدم، فإنه يستجيب بطرق نمطية للغاية. وفي بعض الأحيان، قد يقتبس ملاحظة من تأليف الإنسان أو قصيدة مرتبطة بشيء قاله المستخدم؛ ربما في اليوميات. ولكن الصعوبات التي تواجه معالجة اللغات الطبيعية تنطوي على أن النص المنشأ باستخدام الكمبيوتر لا يُحتمل أن يكون ملائمًا ببراعة. كما أنه قد لا يكون مقبولاً؛ فالمستخدم قد يغضب ويحبّط من رفيق غير قادر على تقديم حتى «ظاهر» الرفقة الحقة. وبالمثل، القط الروبوت الذي يُصدر صوت هرير قد يزعم المستخدم بدلاً من توفير عوامل الرضا النابع من الراحة والاسترخاء.

ثم مرةً أخرى، قد لا يكون الأمر كذلك: «بارو» عبارة عن روبوت على شكل «فقمة صغيرة» تفاعلي ومحبوب، له عيانان سوداوان ورموش جميلة، ويبدو أنه مفيد للعديد من كبار السن والمصابين بالخرف. (في الإصدارات المستقبلية، سيراقب الروبوت العلامات الحيوية، وينبّه مقدمي الرعاية من البشر بناءً على ذلك).

بإمكان بعض رفاق الذكاء الاصطناعي استخدام تعبيرات وجوههم وحملقة أعينهم للاستجابة بطريقة تبدو عاطفية. تمتلك بعض الروبوتات «البشرة» المرنة التي تملو هيئة عضلات الوجه البشري، وتكونها يمكن أن يوحى (للمراقب البشري) بما يصل إلى ١٢ شكلاً من العواطف الأساسية. غالبًا ما تُظهر الأنظمة التي تعمل بالشاشة وجه الشخصية الافتراضية التي تتغير تعبيراتها طبقاً للعواطف التي من المفترض أنها تمرُّ بها. ومع ذلك، كل هذه الأشياء تخاطر بوقوعها ضمن ما يُطلق عليه «الوادي الغريب»؛ فعادةً ما يشعر الناس بعدم الارتياح أو حتى الانزعاج البالغ عند مواجهة مخلوقات شديدة الشبه بالإنسان «ولكنها ليست مماثلة له بالقدر الكافي». لذا يمكن اعتبار الروبوتات — أو الصور الرمزية للشاشة — ذات الوجوه غير البشرية تمامًا بمثابة تهديد.

هل من الأخلاق توفير تلك الرفقة الافتراضية إلى أشخاص يحتاجون إلى العاطفة مسألة محل جدال (انظر الفصل السابع). بالتأكيد بعض الأنظمة التفاعلية بين الكمبيوتر

والإنسان (مثل بارو) يبدو أنها توفر السعادة، بل إنها توفر الرضا الدائم إلى من يعيشون حياة يبدو أنها فارغة. لكن هل هذا كافٍ؟

هناك عمق نظري ضئيل في نماذج «الرفقة». فالأنماط العاطفية لرفاق الذكاء الاصطناعي يجري تطويرها لأغراض تجارية. لا توجد محاولات لجعلها تستخدم عواطفها لحل مشكلاتها الخاصة، ولا لإلقاء الضوء على الدور الذي تلعبه العواطف في عمل العقل ككل. وكأن الباحثين في الذكاء الاصطناعي يرون العواطف على أنها إضافات اختيارية؛ ومن ثم يمكن تجاهلها ما لم تكن حتمية في بعض السياقات الإنسانية الفوضوية.

انتشر هذا الموقف الرافض في الذكاء الاصطناعي حتى وقت قريب نسبياً. عملت روزاليند بيكارد على «الحوسبة الفعّالة»، وأخرجت العواطف من حالة الاستبعاد في أواخر تسعينيات القرن العشرين، ولكن عملها لم يُحلّل العاطفة تحليلًا عميقًا.

من أسباب تجاهل الذكاء الاصطناعي للعاطفة (وملاحظات سيمون نافذة البصيرة بشأنها) هذه المدة الطويلة، هو أن معظم علماء النفس والفلاسفة تجاهلوها هم أيضًا. بعبارة أخرى، لم يفكروا في «الذكاء» على أنه شيء يتطلب العاطفة. وعلى النقيض من ذلك، كان من المفترض أن يعطل التأثير حل المشكلات والعقلانية. وفكرة أن العاطفة يمكن أن تساعد الشخص على اتخاذ القرار بشأن ما يفعله وأفضل طريقة لفعله لم تكن مواكبة للأفكار حينذاك.

لكنها أصبحت بارزة أكثر في النهاية، ويعود الفضل جزئيًا إلى التطورات التي حدثت في علم الأعصاب وعلم النفس الإكلينيكي. ولكن يعود فضل دخولها إلى الذكاء الاصطناعي إلى عالين من علماء الذكاء الاصطناعي، وهما مارفن مينسكي وآرون سلومان، اللذان درسا «العقل ككل» بدلًا من حصر أنفسهما — مثل معظم زملائهم — في زاوية واحدة صغيرة من زوايا العقل.

على سبيل المثال، يركّز مشروع سلومان الجاري «كوج آف» (CogAff) على دور العاطفة في المعمارية الحاسوبية للعقل. وأثر مشروع «كوج آف» في نموذج تعلّم وكيل التوزيع الذكي الخاص بالوعي، والذي انطلق في عام ٢٠١١، ولا يزال يتوسّع (انظر الفصل السادس). والمشروع ألهم أيضًا برنامج «مايندر» MINDER الذي أطلقته مجموعة سلومان في أواخر تسعينيات القرن العشرين.

يحاكي برنامج «مايندر» (الأنماط الوظيفية) للقلق الذي يعتري المربية التي تُترك لرعاية عدة أطفال بمفردها. فهي ليس لديها سوى بضع مهام، وهي: إطعام الأطفال،

ومحاولة حمايتهم من الوقوع في الحفر، واصطحابهم إلى وحدة الإسعافات الأولية إذا وقعوا. وليس لديها سوى بضعة محفزات (أهداف)، وهي: إطعام طفل، ووضع الطفل خلف السياج الواقعي حال توافره، ونقل الطفل من الحفرة لتلقي الإسعافات الأولية، والانتباه للحفر، وبناء سياج، ونقل طفل إلى مسافة آمنة من الحفرة، والتجول حول الحضانة إذا لم يكن هناك محفز نشط في الوقت الحالي.

من ثم، تلك المربية أبسط بكثير من المربية الحقيقية (على الرغم من أنها معقدة أكثر من برنامج التخطيط العادي الذي لديه هدف نهائي واحد). وعلى الرغم من ذلك، فإنها عرضة للاضطرابات العاطفية التي يمكن مقارنتها بأنواع مختلفة من القلق.

المربية المحاكاة عليها أن تستجيب بالطريقة الملائمة إلى الإشارات البصرية التي تتلقاها من بيئتها. بعض تلك الإشارات تحفز (أو تؤثر في) الأهداف الملحة أكثر من غيرها، مثل: طفل يحبو باتجاه الحفرة ويحتاج إلى انتباهها في وقت أقرب من الطفل الجائع، والطفل الذي على وشك الوقوع في الحفرة يحتاج إلى انتباه في وقت أقرب وأقرب. لكن حتى هذه الأهداف التي يمكن تأجيلها قد يلزم التعامل معها في النهاية، وربما تزيد درجة إلحاحها بمرور الوقت. ومن ثم يمكن وضع الطفل الجائع في سريره إذا كان هناك طفل آخر على مشارف الحفرة، ولكن الطفل الذي انتظر فترة أطول من أجل الطعام ينبغي إطعامه قبل الأطفال الذين أطعموا مؤخرًا.

باختصار، يمكن قطع مهمات المربية في بعض الأحيان وتركها أو تأجيلها. ومن ثم لا بد لبرنامج «مايندر» أن يقرر ما هي الأولويات الحالية. يجب اتخاذ تلك القرارات من خلال جلسة، ويمكن أن تؤدي إلى تغييرات متكررة في السلوك. عمليًا، لا يمكن إتمام أي مهمة دون مقاطعة؛ لأن البيئة (الأطفال) تضع كثيرًا من المطالب المتضاربة والمتغيرة باستمرار على النظام. كما هو الحال مع المربية الحقيقية، يزداد القلق ويتدهور الأداء مع زيادة عدد الأطفال؛ حيث إن كل طفل منهم عامل مستقل لا يمكن التنبؤ به. وعلى الرغم من ذلك، فالقلق مفيد، حيث إنه يمكن المربية من إطعام الأطفال بنجاح. بنجاح، ولكن ليس بسلاسة؛ فالهدوء والقلق قطبان متباعدان.

يشير برنامج «مايندر» إلى بعض الطرق التي يمكن للعواطف من خلالها أن تتحكم في السلوك، وتجدول الدوافع المتنافسة بذكاء. لا شك أن المربية البشرية ستواجه عدة أنواع من القلق حسب تغير الموقف. ولكن المهم هنا هو أن العواطف ليست مجرد «مشاعر». بل إنها تتضمن وعيًا وظيفيًا، وكذلك وعيًا بالظواهر (انظر الفصل السادس). وعلى وجه

الذكاء الاصطناعي

التحديد، فإنها آليات حاسوبية تمكّننا من جدولة الأهداف المتنافسة، ولا يمكننا العمل من دونها. (لذا، فإن السيد سبوك عديم المشاعر في «ستار تريك» هو استحالة تطويرية).
إذا أردنا تحقيق الذكاء الاصطناعي العام في أي وقت، فسيتوجب علينا تضمين
العواطف مثل القلق واستخدامها.

الفصل الرابع

الشبكات العصبية الاصطناعية

تتركب الشبكات العصبية الاصطناعية من مجموعة وحدات مترابطة، وكل وحدة قادرة على حوسبة شيء واحد فقط. بذلك الوصف، قد تبدو الشبكات شيئاً مملأً. لكنها تكاد تكون شيئاً سحرياً. ومن المؤكد أنها سحرت الصحفيين. ففي ستينيات القرن العشرين، تحدثت الصحف بحماسة عن شبكات «بيرسيبترون» التي اخترعها فرانك روزنبلات، وهي عبارة عن أجهزة كهروضوئية تعلمت التعرف على الحروف من دون تعليمها بشكل صريح. وعلى وجه التحديد، أحدثت الشبكات العصبية الاصطناعية ضجة في أواسط ثمانينيات القرن العشرين، ولم تنقطع وسائل الإعلام حتى الآن عن الإشادة بها. والضجيج الحالي المرتبط بالشبكات العصبية الاصطناعية له علاقة بالتعلم العميق.

للشبكات العصبية الاصطناعية تطبيقات لا تُحصى، بدايةً من تشغيل سوق الأوراق المالية ورصد تقلبات العملات وحتى التعرف على الكلام والوجوه. ولكن ما يأسر العقل «الطريقة التي تعمل بها» تلك الشبكات.

تعمل حفنة صغيرة من تلك الشبكات على جهاز خاص متوازٍ، أو حتى على مزيج من الأجهزة/الأجهزة العصبية التي تجمع بين الخلايا العصبية الحقيقية ودوائر السيليكون. لكن عادةً ما تُحاكى الشبكة باستخدام جهاز فون نيومان. وهذا يعني أن الشبكات العصبية الاصطناعية عبارة عن أجهزة افتراضية تعمل بالمعالجة المتوازية، وتنفَّذ على أجهزة الكمبيوتر الكلاسيكية (انظر الفصل الأول).

إنها مثيرة للاهتمام جزئياً لأنها تختلف كثيراً عن الأجهزة الافتراضية للذكاء الاصطناعي الرمزي. فالتعليمات المتسلسلة يحل محلّها التوازي الكثيف، والتحكم من أعلى إلى أسفل تحل محلّه المعالجة من أسفل إلى أعلى، والمنطق تحل محلّه الاحتمالية. أما

الجانب الديناميكي المتغير باستمرار في الشبكات العصبية الاصطناعية فيتناقض تناقضاً صارخاً مع البرامج الرمزية.

إضافة إلى ذلك، تمتلك العديد من الشبكات خاصية خارقة، وهي تولّد التنظيم الذاتي من رحم البداية العشوائية. (كذلك شبكات بيرسيبترون التي اخترعت في ستينيات القرن العشرين تمتلك تلك الخاصية، ومن هنا كان انتشارها في الأخبار الصحفية). يبدأ النظام ببنية معمارية عشوائية (أوزان وروابط عشوائية)، ويكيّف نفسه تدريجياً كي ينفذ المهمة المطلوبة.

تمتلك الشبكات العصبية العديد من نقاط القوة، كما أنها أضافت قدرات حاسوبية كبيرة إلى الذكاء الاصطناعي. ومع ذلك، فهي لديها نقاط ضعف. فهي لا يمكنها تحقيق الذكاء الاصطناعي العام كما هو موصوف في الفصل الثاني. على سبيل المثال، على الرغم من أن بعض الشبكات العصبية الاصطناعية يمكنها عمل استدلال أو تفكير منطقي تقريبي، فهي لا تستطيع تمثيل الدقة مثلما يمثلها الذكاء الاصطناعي الرمزي. (س: ما مجموع $2 + 2$ ؟ ج: من المرجح ٤. حقاً؟) كذلك يصعب نمذجة التسلسل الهرمي في الشبكات العصبية الاصطناعية. بعض الشبكات (المتكررة) تستخدم الشبكات التفاعلية لتمثيل التسلسل الهرمي، ولكن بدرجة محدودة.

بفضل الحماسة الحالية للتعليم العميق، فقد أصبحت شبكات الشبكات أقل ندرة الآن عما كانت عليه من قبل. ومع ذلك، فهي لا تزال بسيطة نسبياً. لا بد أن الدماغ البشري يتكون من عدد لا يُحصى من الشبكات على العديد من المستويات، وتتفاعل بطرق معقدة للغاية. باختصار، الذكاء الاصطناعي العام لا يزال بعيداً للغاية.

الآثار الأوسع نطاقاً للشبكات العصبية الاصطناعية

الشبكات العصبية الاصطناعية انتصار للذكاء الاصطناعي الذي يُعتبر من علوم الكمبيوتر. لكن آثارها النظرية أبعد من ذلك بكثير. بسبب بعض أوجه الشبه العامة مع المفاهيم الإنسانية والذاكرة، يهتم علماء الأعصاب وعلماء النفس والفلاسفة بالشبكات العصبية الاصطناعية.

اهتمام علم الأعصاب ليس جديداً. في الحقيقة، لم يقصد روزنبلات من شبكات بيرسيبترون أن تكون مصدرًا للأدوات المفيدة عملياً، ولكن قصد أن تكون «نظرية عصبية

نفسية». وعلى الرغم من أوجه الاختلاف العديدة بين الشبكات والدماغ، أصبحت الشبكات مهمة في علم الأعصاب الحاسوبي.

يهتم علماء النفس أيضًا بالشبكات العصبية الاصطناعية، ولم يتخلف عنهم الفلاسفة كثيرًا. على سبيل المثال، تسبَّب أحد الأمثلة في منتصف ثمانينيات القرن العشرين في إثارة ضجة خارج صفوف الذكاء الاصطناعي الاحترافي. من الواضح أن تلك الشبكة تعلمت استخدام زمن الماضي كثيرًا مثل الأطفال؛ إذ بدأت بعدم ارتكاب أخطاء ثم أفرطت في التنظيم قبل أن تحقق الاستخدام الصحيح لكل من الأفعال المنتظمة وغير المنتظمة؛ ومن ثم الفعل go/went أفسح الطريق للصيغة go/goed. كان ذلك ممكنًا لأن المدخلات التي جرى توفيرها لها تعكس الاحتمالات المتغيرة للكلمات التي يسمعها الطفل عادة؛ لم تكن الشبكة تطبِّق القواعد النحوية الفطرية.

كان هذا مهمًا لأن معظم علماء النفس (والعديد من الفلاسفة) حينذاك قبلوا أقوال نعوم تشومسكي بأن الأطفال يجب أن يعتمدوا على القواعد اللغوية الفطرية حتى يتعلموا النحو، وبأن التنظيم المفرط لدى الأطفال كان دليلًا لا يقبل الجدل على تطبيق تلك القواعد. أثبتت شبكة زمن الماضي عدم صحة أيٍّ من تلك الأقوال. (بالطبع لم تثبت أن الأطفال ليس لديهم قواعد فطرية، بل ببساطة أثبتت أنهم لا يحتاجون إليها).

مثال آخر مثير للاهتمام على نطاق واسع، وهو البحث عن «المسارات التمثيلية»، وهو مستوحى في الأساس من علم نفس النمو. وهنا (وكذلك في التعلم العميق)، تسجَّل البيانات المدخلة التي كانت مربكة في البداية على مستويات تعاقبية، بحيث تُمسك القواعد المنتظمة الأقل وضوحًا بالإضافة إلى القواعد المنتظمة الواضحة. هذا لا يرتبط بنمو الطفل فحسب، بل يرتبط أيضًا بالمناقشات النفسية والفلسفية بشأن التعلم الاستقرائي. وبذلك يتضح أن التوقعات السابقة (البنية الحاسوبية) مطلوبة من أجل تعلم الأنماط في البيانات المدخلة، وأن هناك قيودًا لا مفرَّ منها على الترتيب الذي يجري من خلاله تعلُّم الأنماط المختلفة. باختصار، منهجية الذكاء الاصطناعي هذه مثيرة للاهتمام بعدة طرق، كما أنها بالغة الأهمية من الجانب التجاري.

المعالجة الموزعة المتوازية

توجد فئة من فئات الشبكات العصبية الاصطناعية التي تجذب انتباه كثيرين، ألا وهي التي تنفذ المعالجة الموزعة المتوازية. في الحقيقة، عندما يشير الناس إلى «الشبكات العصبية» أو

«الترابطية» (مصطلح قلَّ استخدامه في هذه الأيام)، فعادةً ما يقصدون المعالجة الموزعة المتوازية.

نظرًا إلى الطريقة التي تعمل بها شبكات المعالجة الموزعة المتوازية، فإنها تتشارك في أربعة نقاط قوى أساسية. ترتبط هذه النقاط بكلٍّ من التطبيقات التكنولوجية وعلم النفس النظري (وبفلسفة العقل أيضًا).

النقطة الأولى هي قدرتها على تعلُّم الأنماط والعلاقات بين تلك الأنماط عن طريق عرض الأمثلة بدلاً من برمجتها برمجة صريحة.

النقطة الثانية هي تقبُّل الأدلة «الفوضوية». تلك الأنظمة يمكنها تحقيق الاكتفاء المقيد، بحيث تُعطي معنىً منطقيًا من الأدلة المتعارضة جزئيًا. إنها لا تتطلب تعريفات دقيقة معبرًا عنها في صورة قوائم من الشروط الضرورية والوافية. بل إنها تتعامل مع المجموعات المتداخلة من التشابهات المتجانسة؛ وهي سمة لمفاهيم البشر كذلك.

نقطة القوة الثالثة هي القدرة على التمييز بين الأنماط التالفة بالكامل والتالفة جزئيًا. فهي تحتوي على ذاكرة قادرة على معالجة المحتوى. وهكذا هم البشر، ولنضرب مثالًا بالتعرف على النغمة من أول بضعة ألحان، أو عندما تتخللها العديد من الأخطاء عند العزف.

نقطة القوة الرابعة هي الدقة. فشبكة المعالجة الموزعة المتوازية التي تفتقد بعض العقد لا تعطي نتائج لا معنى لها ولا تتوقف. إنها تُظهر تدهورًا حميدًا، وفيه يسوء الأداء تدريجيًا مع زيادة التلف. لذا فهي ليست هشّة، مثل البرامج الرمزية.

تلك الفوائد ناتجة عن «التوزيع» في المعالجة الموزعة المتوازية. ليس كل الشبكات العصبية الاصطناعية تتضمن المعالجة الموزعة. في الشبكات المحلية (مثل «ووردنت»؛ انظر الفصل الثاني)، يتم تمثيل المفاهيم بعقد منفردة. في الشبكات الموزعة، يخزّن المفهوم (يوزّع) عبر النظام بأكمله. في بعض الأحيان، تُدمج المعالجة المحلية مع الموزعة، ولكن لا يحدث ذلك كثيرًا. الشبكات المحلية الخالصة ليست شائعة أيضًا؛ لأنها تفتقر إلى نقاط القوة الأساسية في المعالجة الموزعة المتوازية.

يمكن القول إن الشبكات الموزعة عبارة عن شبكات محلية في أساسها؛ لأن كل وحدة تتوافق مع ميزة دقيقة واحدة؛ على سبيل المثال، رقعة لون صغيرة في مكان معين في المجال البصري. ولكن تُعرّف تلك الشبكات على مستوى أقل بكثير من المفاهيم؛ فالمعالجة الموزعة المتوازية تتضمن حوسبة «شبه رمزية». إضافة إلى ذلك، كل وحدة يمكن أن تكون جزءًا من عدة أنماط كلية مختلفة، ومن ثمّ تساهم في العديد من «المعاني» المختلفة.

يوجد العديد من أنواع أنظمة المعالجة الموزعة المتوازية. كل الأنظمة مكوّنة من ثلاث طبقات أو أكثر من الوحدات المتصلة، وكل وحدة لا تقدر إلا على حوسبة شيء واحد بسيط. ولكن الوحدات تختلف.

تنشط الوحدة في طبقة المدخلات متى قُدمت ميزتها الدقيقة إلى الشبكة. تنشط وحدة المخرجات عندما تتحفز بالوحدات المتصلة بها، ونشاطها يجري إيصاله إلى الشخص المستخدم. الوحدات المخفية في الطبقة (الطبقات) الوسطى ليس لها اتصال مباشر بالعالم الخارجي. بعض تلك الوحدات محددة؛ بمعنى أنها تنشط — أو لا تنشط — بناءً على تأثيرات وصلاتها فحسب. بعضها الآخر عشوائي؛ بمعنى أنه يعتمد تنشيطها أو عدمه جزئياً على قدر من التوزيع الاحتمالي.

تختلف الوصلات أيضاً. فبعض الوصلات تكون ذات «تغذية أمامية»، بحيث تمرّ الإشارات من طبقة دنيا إلى طبقة أعلى. وبعضها يرسل الإشارات بـ «التغذية الراجعة» في الاتجاه المعاكس. بعضها «جانبى» يربط الوحدات داخل الطبقة نفسها. وبعضها يجمع بين التغذية الأمامية والتغذية الراجعة كما سنرى. ومثل التشابكات العصبية في الدماغ، تكون الوصلات إما محفزة وإما مثبطة. إنها تتفاوت في القوة، أو «الوزن». يعبر عن الأوزان بالأعداد ما بين +١ إلى -١. كلما زاد وزن الرابط المحفز (أو المثبط)، زاد (أو قل) احتمال أن تنشط الوحدة التي تستقبل الإشارة.

تتضمن المعالجة الموزعة المتوازية تمثيلاً موزعاً؛ لأن كل مفهوم تمثله حالة الشبكة بكاملها. قد يبدو هذا محيراً، بل متناقضاً. بالتأكيد هذا يختلف كثيراً عن طريقة تحديد التمثيلات في الذكاء الاصطناعي الرمزي.

لا يهتم لذلك من يهتمون فقط بالتطبيقات التكنولوجية/التجارية. إذا اقتنعوا بأن أسئلة واضحة بعينها — مثل كيف لشبكة واحدة أن تخرن العديد من المفاهيم أو الأنماط المختلفة — لا تثير مشكلة من الناحية العملية، فإنهم يسعدون بتركها على حالها.

طرح المهتمون أيضاً بالتعقيدات النفسية والفلسفية للذكاء الاصطناعي ذلك «السؤال الواضح». الإجابة هي أن الحالات الإجمالية المحتملة لشبكات المعالجة الموزعة المتوازية كثيرة الأنواع، لدرجة أن القليل منها يتضمن التنشيط المتزامن في انتشار «هذه» الوحدات أو «تلك». ستنتشر الوحدة المفعلّة التنشيط على بعض الوحدات الأخرى فقط. ومع ذلك، تتفاوت تلك «الوحدات الأخرى»؛ بإمكان أي وحدة المساهمة في العديد من أنماط التنشيط

المختلفة. (بوجه عام، تزيد كفاءة التمثيلات «المتفرقة» مع العديد من الوحدات غير النشطة). سيتشبع النظام في النهاية؛ يسأل البحث النظري بشأن الذاكرات الترابطية عن عدد الأنماط التي يمكن أن تخزنها شبكات ذات حجم معين من حيث المبدأ. لكن المهتمين بالأنماط النفسية والفلسفية لن يسعدوا بترك الأمر عند هذا الحد. إنهم مهتمون بمفهوم «التمثيل» نفسه، كما أنهم مهتمون بالنقاشات بشأن ما إذا كان العقل/الدماغ البشري يجري تمثيلات بالفعل أم لا. على سبيل المثال، يجادل أتباع المعالجة الموزعة المتوازية أن ذلك النهج يدحض فرضية نظام الرموز الفيزيائية التي تأصلت في الذكاء الاصطناعي الرمزي، وانتشرت بسرعة في فلسفة العقل (انظر الفصل السادس).

التعلم في الشبكات العصبية

معظم الشبكات العصبية الاصطناعية بإمكانها التعلم. وهذا يتضمن إحداث تغييرات تكيفية في الأوزان وفي الاتصالات أحياناً. عادةً ما يكون تشريح الشبكة ثابتاً، والمقصود بالتشريح هنا عدد الوحدات والروابط التي بينها. وإذا كان الأمر كذلك، فإن التعلم لا يغيّر سوى الأوزان. ولكن في بعض الأحيان يمكن للتعلم — أو التطور (انظر الفصل الخامس) — أن يضيف روابط جديدة ويقطع الروابط القديمة. تستخدم الشبكات البناء تلك الخاصة إلى أقصى حد؛ بمعنى أنها لا تبدأ بوحدات مخفية على الإطلاق، بل إنها تضيفها مع التقدم في التعلم.

بإمكان شبكات المعالجة الموزعة المتوازية التعلم بالعديد من الطرق المختلفة، وقد ضربنا الأمثلة على جميع الأنواع المميزة في الفصل الثاني، وهي: التعلم الموجّه والتعلم غير الموجّه والتعلم المعزّز.

في التعلم الموجّه على سبيل المثال، تتعرف الشبكات على الفئات بعرض أمثلة متنوعة من تلك الفئة، ولا تحتاج أيّ من تلك الشبكات إلى امتلاك كل ميزة «نموذجية». (قد تكون البيانات المدخلة عبارة عن صور بصرية أو أوصاف لفظية أو مجموعات من الأرقام ...) عند تقديم المثال، تستجيب بعض وحدات المدخلات إلى «ميزاتها الدقيقة»، وتنتشر عمليات التنشيط حتى تستقر الشبكة. عندئذٍ، تُقارَن حالة الوحدات المخرجة الناتجة بالنتيجة المرجوة (التي يحددها المستخدم البشري)، ويتم تحفيز المزيد من تغييرات الأوزان (ربما بالانتشار العكسي) بحيث يقل احتمال الأخطاء. بعد العديد من الأمثلة المختلفة اختلافاً طفيفاً، ستكون الشبكة قد طوّرت نمط تنشيط يتطابق مع الحالة النموذجية أو «الأولية»

حتى لو لم تكن قد قابلت حالة مماثلة بالفعل. (وإذا قُدم مثال فاسد الآن بحيث يؤدي إلى تحفيز عدد أقل من وحدات الإدخال ذات الصلة، فلن يكتمل هذا النمط تلقائياً).

يعتمد جزء كبير من تعلم الشبكات العصبية الاصطناعية على قانون «العصبونات التي تنشط معاً ترتبط معاً»، الذي طرحه اختصاصي علم النفس العصبي دونالد هيب في أربعينيات القرن العشرين. يقوِّي قانون هيب للتعلم الوصلات التي تُستخدم كثيراً. عند تنشيط وحدتين في آن واحد، فإن الأوزان تتكيف كي تزيد من احتمال ذلك في المستقبل.

عبر هيب عن قانون «العصبونات التي تنشط معاً ترتبط معاً» بطريقتين، ولكن لم تتسم أيُّ من الطريقتين بالدقة أو المساواة. يعرف الباحثون في الذكاء الاصطناعي اليوم القانون بعدة طرق، وربما يعتمد ذلك على المعادلات المتفاوتة المستمدة من الفيزياء أو نظرية احتمالية بايزي. إنهم يستخدمون التحليل النظري للمقارنة بين الإصدارات المختلفة وتحسينها. ومن ثم قد تكون أبحاث المعالجة الموزعة المتوازية رياضيةً بحتة.

على اعتبار أن شبكة المعالجة الموزعة المتوازية تستخدم قدرًا من قاعدة هيب في التعلم لتكييف أوزانها، فمتى تتوقف؟ الإجابة ليست «عند تحقيق الكمال (إزالة جميع التناقضات)»، بل الإجابة هي «عند تحقيق أقصى قدر من الاتساق».

يحدث «عدم الاتساق» على سبيل المثال عندما ترسل الوحدات ذات الصلة إشارة إلى ميزتين دقيقتين لا تُوجدان معاً في العادة. بإمكان العديد من برامج الذكاء الاصطناعي الرمزي تقييد الاكتفاء، بحيث تقترب من الحل عن طريق إزالة التناقضات بين الألة في أثناء عملية التعلم. لكنها لا تعامل عدم الاتساق باعتباره جزءاً من الحل. فأنظمة المعالجة الموزعة المتوازية مختلفة. ومع ظهور نقاط قوة المعالجة الموزعة المتوازية المدرجة فيما سبق، فيمكنها النجاح في الأداء حتى مع وجود التناقضات. وعندئذٍ يصبح «حلها» هو الحالة الكلية للشبكة عند الحد من التناقضات وليس عند تعطيلها.

من طرق تحقيق ذلك استعارة فكرة «نقطة التوازن» من الديناميكا الحرارية. يعبر عن مستويات الطاقة باستخدام الأعداد، مثلما هي الحال مع الأوزان في المعالجة الموزعة المتوازية. إذا كانت قاعدة التعلم تضاهي قوانين الفيزياء (وإذا كانت الوحدات المخفية عشوائية)، فإن معادلات بولتزمان الإحصائية يمكنها أن تصف التغيرات في الحالتين كليهما.

بإمكان المعالجة الموزعة المتوازية أن تستعير الطريقة المستخدمة في تبريد المعادن بسرعة ولكن بالتساوي. يبدأ التلدين عند درجة حرارة مرتفعة ويبرد تدريجياً. يستخدم

الباحثون في المعالجة الموزعة المتوازنة «خوارزمية محاكاة التلدين»، حيثما كانت تغييرات الوزن في الدورات الأولى القليلة من التوازن أكبر بكثير من التغييرات في الدورات اللاحقة. هذا يُمْكِن الشبكة من تفادي المواقف (الحدود الدنيا المحلية)، حيث يتحقق الاتساق الكلي بالنسبة إلى ما جرى قبل ذلك، ولكن يمكن الوصول إلى مستوى أكبر من الاتساق (ومستوى استقرار أعلى في نقطة التوازن) إذا اضطرب النظام. يمكن ضرب المثل بهزّ كيس من الكرات الزجاجية لإخراج أي كرات مستقرة في الحافة الداخلية؛ بمعنى أنه في البداية ينبغي الهز بقوة ثم ينتهي بالهز بلطف.

توجد طريقة أسرع ومنتشرة الاستخدام أكثر لتحقيق أقصى درجة اتساق، ألا وهي توظيف الانتشار العكسي. ولكن أيّاً كانت القاعدة المطبقة من قواعد التعلم العديدة، فإن حالة الشبكة بالكامل (ولا سيما وحدات المخرجات) عند نقطة التوازن تُعتبر تمثيلاً للمفهوم المعني.

الانتشار العكسي والدماغ والتعلم العميق

يجادل المتحمسون للمعالجة الموزعة المتوازنة أن شبكاتهم واقعية من الناحية البيولوجية أكثر من الذكاء الاصطناعي الرمزي. وصحيح أن المعالجة الموزعة المتوازنة مستلهمة من الدماغ، وأن بعض علماء الأعصاب يستخدمونها لنمذجة الوظائف العصبية. ولكن الشبكات العصبية الاصطناعية مختلفة إلى حد كبير عما يقبع داخل رؤوسنا.

يتمثل أحد الفروق بين (معظم) الشبكات العصبية الاصطناعية والدماغ في الانتشار العكسي. إنها قاعدة تعلم — أو بالأحرى فئة عامة من قواعد التعلم — تُستخدم كثيراً في المعالجة الموزعة المتوازنة. لقد توقّعها بول ويربوس عام ١٩٧٤، وقد وضع جيفري هينتون لها تعريفاً عملياً أكثر في أوائل ثمانينيات القرن العشرين. إنها حل مسألة «إحالة الاستحقاق».

تنشأ تلك المسألة عبر كل أنواع الذكاء الاصطناعي، لا سيما عندما يتغير النظام باستمرار. بالنظر إلى نظام ذكاء اصطناعي معقد ولكنه ناجح، فما هي أكثر الأجزاء مسئولية عن هذا النجاح؟ في الذكاء الاصطناعي التطوري، غالباً ما يجري تعيين الاستحقاق باستخدام خوارزمية «لواء الدلو» (انظر الفصل الخامس). في أنظمة المعالجة الموزعة المتوازنة ذات الوحدات المحددة (غير العشوائية)، فالاستحقاق عادةً ما يعيّن الانتشار العكسي.

تتعقب خوارزمية الانتشار العكسي المسئولية بترتيب عكسي من طبقة المخرجات إلى الطبقات المخفية، وتحدد الوحدات الفردية التي يلزم تكيفها. (يجري تحديث الأوزان لتقليل أخطاء التوقع). تحتاج الخوارزمية إلى معرفة الحالة الدقيقة لطبقة المخرجات عندما تعطي الشبكة الإجابة الصحيحة. (وبذلك يكون الانتشار العكسي هو التعلم الموجّه). تُعقد مقارنات فردية الوحدات بين تلك المخرجات النموذجية والمخرجات التي يجري الحصول عليها بالفعل من الشبكة. وأي فرق بين نشاط وحدة المخرجات في الحالتين يُعتبر خطأً.

تفترض الخوارزمية أن ذلك الخطأ في وحدة المخرجات ناتج عن خطأ (أخطاء) في الوحدات المتصلة بها. بالعمل باتجاه عكسي عبر النظام، يُنسب قدر معين من الخطأ إلى كل وحدة في الطبقة المخفية الأولى بناءً على وزن الاتصال بينها وبين وحدة المخرجات. تتم مشاركة المسئولية عن الخطأ بين كل الوحدات المخفية المتصلة بوحدة المخرجات التي وقع فيها الخطأ. (إذا كانت وحدة مخفية مرتبطة بالعديد من وحدات المخرجات، يجري تجميع المسئوليات المصغرة). عندئذٍ، تُعزى تغييرات الوزن النسبي إلى الوصلات بين الطبقة المخفية والطبقة السابقة.

قد تكون تلك الطبقة طبقةً أخرى (وأخرى ...) من الوحدات المخفية. ولكن في النهاية، ستصبح طبقة المدخلات، وستتوقف التغييرات في الأوزان. تُكرر تلك العملية حتى تُقلل التناقضات في طبقة المخرجات.

على مدار سنوات عديدة، لم يكن الانتشار العكسي يُستخدم إلا في الشبكات التي تحتوي على طبقة مخفية واحدة. كانت الشبكات ذات الطبقات المتعددة نادرة؛ فإنه يصعب تحليلها، بل يصعب استخدامها في التجربة. لكن في الآونة الأخيرة، أحدثت ثورة هائلة — وبعض الضجيج غير المسئول — مع ظهور التعلم العميق. وفيها، يتعلم النظام بنية تصل إلى عمق النطاق بدلاً من مجرد الأنماط السطحية. بعبارة أخرى، يكتشف النظام تمثيل معرفة له مستويات متعددة وليس مستوى واحدًا.

التعلم العميق مثير للاهتمام؛ لأنه يبشّر بتمكين الشبكات العصبية الاصطناعية من التعامل مع التسلسل الهرمي على الأقل. منذ أوائل ثمانينيات القرن العشرين، كافح أتباع الترابطية أمثال هينتون وجيف إلمان من أجل تمثيل التسلسل الهرمي، إما بالجمع بين التمثيل المحلي/الموزّع وإما بتحديد الشبكات المتكررة. (في الواقع، الشبكات المتكررة تعمل باعتبارها سلسلة من الخطوات المنفصلة. باستخدام التعلم العميق، تستطيع الإصدارات

الحديث في بعض الأحيان أن تتنبأ بالكلمة التالية في جملة أو حتى بـ «الفكرة» التالية في فقرة). ولكنها حققت نجاحًا محدودًا (لا تزال الشبكات العصبية الاصطناعية غير مناسبة لتمثيل التسلسلات الهرمية ذات التحديد الدقيق أو المنطق الاستنتاجي).

انطلق التعلم العميق هو الآخر في ثمانينيات القرن العشرين (على يد يورجن شميدهور). ولكن تفجر المجال أكثر في الفترة الأخيرة عندما قدم هينتون طريقة فعّالة في تمكين الشبكات المتعددة الطبقات من اكتشاف العلاقة بين المستويات المتعددة. تتشكل أنظمة التعلم العميق من آلات بولتزمان «المقيدة» (من دون اتصالات جانبية) على ست طبقات. أولاً، تنفذ الطبقات التعلم غير الموجه. يجري تدريب الطبقات واحدة بواحدة باستخدام خوارزمية محاكاة التلدين. تُستخدم مخرجات طبقة باعتبارها مدخلات للطبقة التالية. وعندما تستقر الطبقة الأخيرة، يُضبط النظام بأكمله ضبطًا دقيقًا باستخدام الانتشار العكسي حتى يصل إلى كل المستويات لتعيين الاستحقاق على النحو الملائم.

هذا النهج في التعلم يثير اهتمام اختصاصيي علم الأعصاب المعرفي، وكذلك مهندسو تكنولوجيا الذكاء الاصطناعي. يعود السبب في ذلك إلى «النماذج التوليدية» التي تتعلم التنبؤ (بأرجح) الأسباب وراء المدخلات؛ ومن ثم توفر نموذجًا أسماه هلمهولتز «الإدراك في صورة استدلال لا واعي» عام ١٨٦٧. أي إن الإدراك ليس مسألة تلقي مدخلات بشكل سلبي من أعضاء الحواس. بل إنه يتضمن التفسير النشط، وحتى التنبؤ الاستباقي لتلك المدخلات. باختصار، العين/الدماغ ليس كاميرا.

التحق هينتون بشركة جوجل عام ٢٠١٣؛ ومن ثم سَيُستخدم الانتشار العكسي كثيرًا. تستخدم جوجل بالفعل التعلم العميق في العديد من التطبيقات، ومنها التعرف على الكلام ومعالجة الصور. وفي عام ٢٠١٤، اشترت شركة ديب مايند وقد أتقنت خوارزمية شبكة كيو العميقة التي تستخدمها ألعاب الأتاري الكلاسيكية عن طريق الجمع بين التعلم العميق والتعلم المعزّن، وحاز برنامجها «ألفاجو» الريادة على مستوى العالم عام ٢٠١٦ (انظر الفصل الثاني). تفضّل شركة «آي بي إم» أيضًا التعلم العميق؛ فهي تستخدمه في برامج «واتسون»، ويتم استعارته للعديد من التطبيقات المتخصصة (انظر الفصل الثالث).

لكن إذا كان التعلم العميق مفيدًا بلا شك، فهذا لا يعني أنه مفهوم جيدًا. يُوضح العديد من قواعد التعلم المتعددة الطبقات من الناحية العملية، ولكن التحليل النظري مُربك.

من بين الأسئلة الكثيرة التي لا يوجد إجابة عليها السؤال الآتي: هل يوجد عمق كافٍ لتحقيق أداء يكاد يضاهي أداء العقل البشري؟ (وحدة وجه القطة المذكورة في الفصل الثاني ناتجة عن نظام مكوّن من تسع طبقات). الجهاز البصري لدى الإنسان على سبيل المثال يحتوي على سبعة مستويات تشريحية، لكن كم عدد المستويات التي تضيفها العمليات الحاسوبية في القشرة الدماغية؟ وبما أن الشبكات العصبية الاصطناعية مستوحاة من الدماغ (نقطة لا ينفك التأكيد عليها في الحديث عن التعلم العميق)، فهذا السؤال طبيعي. ولكن المسألة ليست وثيقة الصلة كما يبدو.

الانتشار العكسي انتصار حاسوبي. ولكنه غير بيولوجي إلى أقصى حد. لا يمكن أن تنتج «خلية جدة» لوجه القطة في الدماغ (انظر الفصل الثاني) عن عمليات مثل التي تحدث في التعلم العميق. التشابكات العصبية الحقيقية تتسم بالتغذية الأمامية الخالصة؛ بمعنى أنها لا تنقل في الاتجاهين كليهما. يحتوي الدماغ على اتصالات التغذية الراجعة في اتجاهات متعددة، ولكن كل اتجاه يسير في مسار واحد فقط. هذا مجرد فرق واحد من بين العديد من الفروق بين الشبكات العصبية الحقيقية والاصطناعية. (هناك فرق آخر، وهو أن الشبكات في الدماغ ليست منظمة تنظيمًا هرميًا صارمًا، على الرغم من أن الجهاز البصري غالبًا ما يوصف بتلك الطريقة).

حقيقة أن الدماغ يحتوي على كلٍّ من الوصلات ذات التغذية الأمامية والراجعة ضرورية من أجل نماذج «الترميز التنبؤي» الخاص بالتحكم الحسي الحركي، وهو ما يسبب إثارة عظيمة في علم الأعصاب. (هذه النماذج أيضًا قائمة على عمل هينتون). ترسل مستويات الأعصاب العليا رسائل إلى الطبقات الأدنى بحيث تتنبأ بالإشارات الواردة من المستشعرات، ولا يرسل إلى الطبقات الأعلى سوى رسائل «الخطأ» غير المتوقعة. تكرار الدورات لهذا النوع يولف الشبكات التنبؤية؛ ومن ثم تتعلم ما ينبغي توقعه بالتدرج. يتحدث الباحثون عن «دماغ بايزي»؛ لأن التنبؤات يمكن تفسيرها من منظور إحصائيات بايزي (انظر الفصل الثاني)، كما أن نماذج الكمبيوتر قائمة على تلك الإحصائيات في الحقيقة.

بالمقارنة مع الدماغ، فإن الشبكات العصبية الاصطناعية بالغة التنظيم والبساطة وقلة العدد وجذب المعلومات. إنها بالغة التنظيم لأن الشبكات التي يبنها الإنسان تعطي الأولوية للذكاء والقوة الرياضية، ولكن الدماغ المتطور بيولوجيًا لا يفعل ذلك. وبالغة البساطة لأن خلية عصبية واحدة معقدة حسابيًا بقدر تعقيد نظام معالجة موزعة متوازية

أو حتى جهاز كمبيوتر صغير، ويوجد ٣٠ نوعًا مختلفًا من الخلايا العصبية. وبالغة قلة العدد لأنه حتى الشبكات العصبية الاصطناعية التي تحتوي على ملايين الوحدات البالغة الصغر مقارنةً بالدماغ البشري (انظر الفصل السابع). وبالغة جذب المعلومات لأن الباحثين في الشبكات العصبية الاصطناعية لا يتجاهلون العوامل الزمنية مثل الترددات العصبية التصاعدية والمزامنات فحسب، بل يتجاهلون الفيزياء الحيوية للعمود الفقري الشجيري والمعدلات العصبية والتيارات الاتصالية العصبية ومرور الأيونات.

كل واحد من مواطن الضعف تلك آخذ في التضاؤل. زيادة قوة أجهزة الكمبيوتر يُمكِّن الشبكات العصبية الاصطناعية من تكوين المزيد من الوحدات الفردية. تُبنى نماذج ذات تفاصيل أكثر بكثير من الخلايا العصبية المفردة، والتي تعالج الوظائف الحاسوبية لكل العوامل العصبية التي ذكرناها للتو. حتى الجذب يتضاءل في الحقيقة كما يتضاءل في المحاكاة (تجمع بعض الأبحاث «العصبية» بين الخلايا العصبية الحية والرقاقات المصنوعة من السيليكون). وبقدر ما تحاكي خوارزمية شبكة كيو العميقة العمليات في القشرة البصرية والحسين (انظر الفصل الثاني)، لا شك أن الشبكات العصبية الاصطناعية ستستعير وظائف أخرى من علم الأعصاب.

وعلى الرغم من ذلك، يظل صحيحًا أن الشبكات العصبية الاصطناعية لا تشبه الدماغ من نواحٍ مهمة لا حصر لها، وبعضها لا نعرفه حتى الآن.

إخفاق الشبكة

تعود الإثارة بشأن المعالجة الموزعة المتوازية بدرجة كبيرة إلى حقيقة أن الشبكات العصبية الاصطناعية (المعروفة أيضًا باسم الترابطية) قد وصلت إلى طريق مسدود قبل ذلك بعشرين عامًا. كما أُشير في الفصل الأول، وردَّ هذا الحكم في نقد لاذع من مارفن مينسكي، وسيمور بابرت في ستينيات القرن العشرين، وكلاهما له سمعة رائنة في مجتمع الذكاء الاصطناعي. بحلول ثمانينيات القرن العشرين، بدا أن الشبكات العصبية الاصطناعية لم تصل إلى طريق مسدود فحسب، بل وصلت إلى نهايتها. وفي الحقيقة، هُملت السبرانية بوجه عام (انظر الفصل الأول). بل انتقلت كل التمويلات البحثية تقريبًا إلى الذكاء الاصطناعي الرمزي.

بدأت بعض أنواع الشبكات العصبية الاصطناعية الأولى مبشرةً إلى حد بعيد. يمكن لشبكات «بيرسيبترون» الذاتية التنظيم التي طُوِّرها روزنبلات — التي كثيرًا ما يرصدها

الصحفيون المفتونون بها — أن تتعلم التعرف على الأنماط بالرغم من أنها بدأت من حالة عشوائية. وقد أدلى روزنبلات بمزاعم طموحة للغاية عن إمكانيات نهجه، بحيث تناولت كل الجوانب النفسية لدى الإنسان. ومن باب التأكيد، أشار إلى تقييدات معينة. لكن «برهان التقارب» المثير للجدل كان قد ضمن تعلم شبكات «بيرسيبترون» أي شيء يمكن برمجتها على إنجازها. وتلك من نقاط القوة.

لكن في أواخر ستينيات القرن العشرين، قدّم كلٌّ من مينسكي وبابرت براهينهما. استخدمتا الرياضيات، وأظهرا أن شبكات «بيرسيبترون» البسيطة لا يمكنها إنجاز أشياء معينة؛ إذ يتوقع المرء بحدسه أن تلك الأنظمة قادرة على إنجازها (برغم أنه بمقدور الذكاء الاصطناعي التقليدي الجميل أن ينجز تلك الأشياء بسهولة). ومثل نظرية التقارب التي وضعها روزنبلات، فإن براهينهما لا تنطبق إلا على الشبكات الأحادية الطبقات. ولكن كان «حكمهما الأولي» أن الأنظمة المتعددة الطبقات يمكن أن يغلبها الانفجار التوافقي. بعبارة أخرى، لن تتوسع شبكات «بيرسيبترون».

اقتنع أغلب علماء الذكاء الاصطناعي أن الترابطية لن تنجح. لم يأبه عدد قليل بذلك، وأجروا أبحاثًا عن الشبكات العصبية الاصطناعية. في الواقع، تحقق تقدم كبير بشأن تحليل الذاكرة الترابطية (على يد كريستوفر لونجيت هيجنز وديفيد ويلشاو، وأخيرًا على يد جيمس أندرسون وتيوفو كوهونين وجون هوبفيلد). ولكن ظل هذا العمل طي الخفاء. لم تعرّف المجموعات المعنية نفسها بأنهم باحثون في الذكاء الاصطناعي، وتجاهلهم الذين عرّفوا أنفسهم بتلك الصفة بوجه عام.

بدّد وصول المعالجة الموزعة المتوازية تلك الشكوك. فبالإضافة إلى بعض النماذج الوظيفية (مثل أداة تعلم الزمن الماضي)، كانت هناك نظريتان جديدتان للتقارب؛ الأولى تضمن أن نظام المعالجة الموزعة المتوازية القائم على معادلات بولتزمان في الديناميكا الحرارية سيصل إلى نقطة التوازن (على الرغم من أنه قد يصل بعد مدة طويلة)، والثانية تثبت أن الشبكة الثلاثية الطبقات يمكنها حل أي مسألة تُعرض عليها من حيث المبدأ. (تحذير صحي: كما هي الحال في الذكاء الاصطناعي الرمزي، غالبًا ما يكون أصعب جزء في التدريب هو تقديم مسألة بطريقة يمكن أن تندرج ضمن المدخلات في جهاز الكمبيوتر). وبطبيعة الحال، اندلعت النقاشات المثيرة. ومن ثم تبعثر إجماع الآراء بشأن الذكاء الاصطناعي السائد.

افترض الذكاء الاصطناعي الرمزي أن التفكير الحدسي السلس يشبه الاستدلال الواعي تمامًا، ولكن من دون الوعي. وكان الباحثون في المعالجة الموزعة المتوازية يقولون

إن هذه أنواع تفكير مختلفة من حيث الجوهر. كل رواد حركة المعالجة الموزعة المتوازية (ديفيد روميلهارت وجاي ماكلياند ودونالد نورمان وهينتون) أشاروا إلى أن كلا النوعين من الأساسيات في علم النفس الإنساني. ولكن الدعاية عن المعالجة الموزعة المتوازية — وتفاعل عامة الناس تجاهها — انطوت على أن الذكاء الاصطناعي الرمزي مضيعة للوقت رغم أنه يُعتبر دراسة للعقل. لكن تغيّر الوضع تمامًا.

كذلك الممول الأساسي للذكاء الاصطناعي — وزارة الدفاع الأمريكية — تراجعت عن موقفها. بعد اجتماع طارئ عام ١٩٨٨، اعترفت الوزارة أن تجاهلها السابق للشبكات العصبية الاصطناعية لم يكن في محله. ومن بعدها، أُغِدَّت الأموال على أبحاث المعالجة الموزعة المتوازية.

بالنسبة إلى مينسكي وبابرت، فهما لم يغيرا رأيهما. في الإصدار الثاني من كتابهما المناهض للشبكات العصبية الاصطناعية، قالا «إن مستقبل تعلم الآلة القائم على الشبكات [ثري] إلى أبعد الحدود». ولكنهما أصرّا على أن الذكاء العالي المستوى لا يمكن أن ينشأ من العشوائية الخالصة، ولا من نظام غير تسلسلي بالكامل. وعليه، لا بد أن يعمل الدماغ في بعض الأحيان مثل معالج تسلسلي، وسيُضطر الذكاء الاصطناعي على المستوى البشري أن يوظف أنظمة مختلطة. تعللًا بأن نقدهما لم يكن العامل الوحيد الذي أدى إلى سنوات الجذب بشأن الشبكات العصبية الاصطناعية؛ وما كان ذلك إلا لسبب واحد، وهو عدم توافر القوة الكافية لأجهزة الكمبيوتر. كما أنهما أنكرا محاولة تحويل أموال الأبحاث إلى الذكاء الاصطناعي الرمزي. بعبارة أخرى، «لم نفكر في عملنا وكأنه محاولة لقتل «سنوايت»، بل اعتبرناه طريقة لفهمها».

كانت هذه حججًا علمية معتبرة. ولكن نقدهم الأولي كان لاذعًا. (كانت المسودة لاذعة أكثر، ونُصحهم الزملاء الودودين بتخفيف حدتها وإبراز النقاط العلمية أكثر). ولا عجب من أن هذه النصيحة أثارت العاطفة. استاء مناصرو الشبكات العصبية الاصطناعية من أعماقهم بسبب عدم رؤية مشروعهم الثقافي المنشأ حديثًا. بل أحدثت المعالجة الموزعة المتوازية ضجة أكبر. تضمن أقول الشبكات العصبية الاصطناعية الغير والنكاية والتعظيم الذاتي والشماتة المرحلة: «أخبرناك هذا من قبل!»

كانت هذه الحلقة مثالًا بارزًا على إخفاق علمي، وهي ليست الوحيدة التي ظهرت في مجال الذكاء الاصطناعي. أُقحمت الخلافات النظرية في العواطف الشخصية والمنافسات، وكان التفكير النزيه نادرًا. انتشرت الانتقادات اللاذعة والأخبار الصحفية أيضًا. الذكاء الاصطناعي ليس مسألة تخلو من العاطفة.

الوصلات ليست كل شيء

تقول معظم حسابات الشبكات العصبية الاصطناعية إن الشيء المهم الوحيد بشأن الشبكة العصبية هو بنيتها التشريحية. ما الوحدات المتصلة بوحدات أخرى، وما مدى قوة الأوزان؟ لا شك أن هذه الأسئلة بالغة الأهمية. ومع ذلك، أظهر علم الأعصاب الحديث أن الدوائر البيولوجية يمكن أن تغير وظيفتها الحاسوبية في بعض الأحيان (ولا تزيد أو تقلل من درجة احتماليتها فحسب)، والسبب في ذلك أن المواد الكيميائية تتغلغل في الدماغ.

أكسيد النيتروز على سبيل المثال ينتشر في كل الاتجاهات، ويستمر تأثيره حتى يتحلل، وتعتمد قوة التأثير على نسبة التركيز في النقاط ذات الصلة. (قد يتفاوت معدل التحلل حسب الإنزيمات). ومن ثم يعمل أكسيد النيتروز على جميع الخلايا داخل مساحة معينة من القشرة، سواء كانت متصلة اتصالاً متشابكاً أم لا. تختلف الديناميكيات الوظيفية للأنظمة العصبية المعنية اختلافاً كبيراً عن الشبكات العصبية الاصطناعية «الخالصة»؛ لأن إشارات الحجم تحل محل الإشارات من نقطة إلى نقطة. عُثر على تأثيرات مماثلة لأول أكسيد الكربون وكبريتيد الهيدروجين، والجزيئات المعقدة مثل السيروتونين والدوبامين.

قد يقول مشكك في الذكاء الاصطناعي: «ما أكثر ما قيل في الشبكات العصبية الاصطناعية! لا توجد مواد كيميائية داخل أجهزة الكمبيوتر!» وربما يضيف: «وقد يترتب على ذلك عدم قدرة الذكاء الاصطناعي على نمذجة الأمزجة أو العواطف؛ إذ إنهما يعتمدان على الهرمونات والمعدلات (المهيات) العصبية.» ذلك الاعتراض نفسه قاله عالم النفس أولريك نيسر في أوائل ستينيات القرن العشرين، وقاله الفيلسوف جون هاوجلاند بعد عدة سنوات في نقده المؤثر عن «مذهب الإدراكية». يقولان إن الذكاء الاصطناعي يمكن أن يعد نماذج للتفكير ولكن دون أي تأثير.

ومع ذلك، هذه الاكتشافات العلمية العصبية ألهمت بعض الباحثين في الذكاء الاصطناعي لتصميم شبكات عصبية اصطناعية من نوع جديد تماماً، حيث لا تكون الوصلات هي كل شيء. في شبكات «جاس نت» (GasNet)، بإمكان بعض العقد المبعثرة عبر الشبكة أن تطلق «غازات» يجري محاكاتها. هذه العقد قابلة للانتشار، وتعديل الخصائص الجوهرية للعقد والتوصيلات الأخرى بعدة طرق، بناءً على التركيز. مقدار حجم الانتشار مهم، وكذلك شكل المصدر (إنشاء نموذج على شكل كرة مجوفة وليس مصدرًا نقطيًا). ومن ثم، العقدة الواحدة سيختلف سلوكها مع اختلاف الوقت. وفي حالات غازية معينة، ستؤثر عقدة في أخرى على الرغم من أنه لا يوجد رابط مباشر. أهم ما في

المسألة التفاعل بين الغاز والتوصيلات الكهربائية داخل النظام. وبما أن الغاز لا ينبعث إلا في حالات معينة، كما أنه ينتشر ويتحلل بنسب متفاوتة، فإن هذا التفاعل معقد ديناميكياً. استخدمت تكنولوجيا «جاس نت» على سبيل المثال في تصميم «أدمغة» للروبوتات المستقلة. وجد الباحثون أن سلوكاً معيناً قد يتضمن شبكتين فرعيتين «غير متصلتين»، ولكنهما تعملان معاً بسبب التأثيرات التعديلية. وجدوا أيضاً «مستكشف اتجاهات» قادراً على استخدام مثلث من الورق المقوى، واتخاذ أداة مساعدة على الملاحظة، ويمكن أن تتطور في شكل شبكات فرعية غير متصلة جزئياً. سبق أن طُوروا شبكة متصلة بالكامل لفعل ذلك (انظر الفصل الخامس)، لكن الإصدار ذا التعديل العصبي تطور سريعاً، وامتاز بكفاءة أكبر.

لذا عدل بعض الباحثين في الشبكات العصبية الاصطناعية عن دراسة تشريح (الوصلات) إلى التعرف على الكيمياء العصبية كذلك. والآن، يمكن محاكاة قواعد التعلم المختلفة وتفاعلاتها الزمنية بأخذ التعديل العصبي بعين الاعتبار. التعديل العصبي ظاهرة تماثلية، وليست رقمية. كذلك يجب ألا تتوقف تركيزات الجزيئات المنتشرة عن التغيير. والأكثر من ذلك، يعكف الباحثون في الذكاء الاصطناعي (باستخدام «رقاقات التكامل الشاسع النطاق» (VLSI)) على تصميم شبكات تجمع بين الوظائف التماثلية والرقمية. تجري نمذجة الميزات التماثلية بناءً على تشريح الخلايا العصبية البيولوجية وبنيتها الفسيولوجية، بما في ذلك مسارات الأيونات عبر غشاء الخلية. تُستخدم هذه الحوسبة «العصبية» على سبيل المثال لمحاكاة أنماط الإدراك والتحكم الحركي. يخطط بعض الباحثين في الذكاء الاصطناعي لاستخدام الحوسبة العصبية داخل نمذجة «الدماغ الكامل» (انظر الفصل السابع).

ذهب آخرون إلى ما هو أبعد من ذلك؛ فبدلاً من نمذجة الشبكات العصبية الاصطناعية بالسيليكون بالكامل، فإنهم يبنون (أو يطورون، انظر الفصل الخامس) شبكات تتكون من أقطاب كهربية مصغرة وخلايا عصبية حقيقية. على سبيل المثال، عند محاكاة القطب «إكس» والقطب «واي» اصطناعياً، فإن النشاط الناتج في الشبكة «الرطبة» يؤدي إلى تحفيز بعض الأقطاب الأخرى مثل القطب «زد»؛ ومن ثم تنفيذ «بوابة اقتران». هذا النوع من الحوسبة في طور المهد (وقد تصوّره دونالد ماكاي في أربعينيات القرن العشرين). ولكنه قد يكون مثيراً.

الأنظمة المختلطة

مفهوم أنه يمكن وصف الشبكات التماثلية/الرقمية وشبكات الأجهزة/العصبية المذكورة آنفاً بأنها أنظمة «مختلطة». ولكن عادةً ما يُستخدم هذا المصطلح للإشارة إلى برامج الذكاء الاصطناعي التي تشمل كلاً من معالجة المعلومات الرمزية والترابطية.

هذا ما وصفه مينسكي في بيانه عام ١٩٥٦ بأنه يمثل ضرورة على الأرجح، وبالفعل جمعت بضعة برامج رمزية أولية بين المعالجة المتوازية والتسلسلية. ولكن هذه المحاولات كانت نادرة. وكما رأينا، أردف مينسكي حتى أوصى بالأنظمة المختلطة الرمزية/الشبكات العصبية الاصطناعية بعد وصول المعالجة الموزعة المتوازية. لكن لم تتوال هذه الأنظمة من فورها (على الرغم من أن هينتون بنى شبكات تجمع بين الاتصالات المحلية والموزعة، والهدف هو تمثيل تسلسلات هرمية جزئية/كاملة مثل أشجار العائلة).

في الحقيقة، لا يزال التكامل بين المعالجة في الشبكات الرمزية والشبكات العصبية غير شائع. المنهجيتان — المنطقية والاحتمالية — مختلفتان تماماً، لدرجة أن معظم الباحثين لديهم خبرة في منهجية واحدة فقط.

على الرغم من ذلك، طُورت بعض الأنظمة المختلطة بالفعل، وهنا يجري تمرير التحكم بين الوحدات الرمزية ووحدات المعالجة الموزعة المتوازية حسب الاقتضاء. لذلك، يعتمد النموذج على نقاط القوة التكميلية لكلا النهجين.

من بين الأمثلة على ذلك خوارزميات لعب الأتاري التي طوّرتها شركة ديب مايند (انظر الفصل الثاني). تجمع هذه الخوارزميات بين التعلم العميق والذكاء الاصطناعي التقليدي الجميل لتعلم كيفية ممارسة مجموعة ألعاب متنوعة من ألعاب الكمبيوتر عن طريق التعلم البصري. إنها تستخدم التعلم المعزّز؛ لا توجد قواعد يدوية، ولا يوجد سوى وحدات البيكسل والنتائج الرقمية في كل خطوة. يُنظر في العديد من القواعد/الخطط في آن واحد، والقاعدة المبشرة أكثر تقرر الإجراء التالي. (ستركز الإصدارات المستقبلية على الألعاب الثلاثية الأبعاد مثل «ماين كرافت»، وعلى التطبيقات مثل السيارات دون سائق). تشمل الأمثلة الأخرى أنظمة العقل الكامل «التحكم المتكيف مع التفكير والعقلانية*» و«التعلم الاتصالي المزود بتحفيز القاعدة التكيفية عبر الإنترنت» (انظر الفصل الثاني) ونظام «تعلّم وكيل التوزيع الذكي» (انظر الفصل السادس). دُرست هذه الأنظمة دراسة عميقة من حيث علم النفس المعرفي؛ لأنها طُورت لأغراض علمية لا لأغراض تكنولوجية.

تأخذ بعض النماذج المختلطة في اعتبارها أنماطاً محددة من علم الأعصاب أيضاً. على سبيل المثال، نشر عالم الأعصاب الإكلينيكي تيموثي شاليس — بالتعاون مع رائد المعالجة الموزعة المتوازية نورمان — نظرية مختلطة عن الإجراء المعتاد (فرط التعلم) عام ١٩٨٠، وقد نُفذت فيما بعد. تشرح النظرية أخطاءً شائعة معينة. على سبيل المثال، كثيراً ما ينسى المصابون بالسكتة الدماغية أن الخطاب ينبغي أن يوضع في المظروف قبل لعق اللسان اللاصق، أو قد يخلدون إلى النوم عندما يصعدون إلى الطابق العلوي لتغيير ملابسهم، أو قد يأخذون الغلاية بدلاً من إبريق الشاي. تقع تلك الأخطاء مثل الترتيب والأخذ واستبدال الأشياء من وقت لآخر.

لكن لماذا؟ ولماذا يتعرض المصابون بتلف في الدماغ على وجه الخصوص لتلك الأخطاء؟ تفيد نظرية شاليس الحاسوبية بأن الإجراء المألوف ينتج عن نوعين من التحكم، يمكن أن يتوقفاً أو يُهيئاً في أوقات معينة. النوع الأول تلقائي، وهو «جدولة التنافس». إنه يتضمن المنافسة (غير الواعية) بين مختلف مخططات العمل ذات التنظيم الهرمي. وينتقل التحكم إلى مخطط تجاوز تنشيطه حدًا معينًا. النوع الآخر (التنفيذي) من آلية التحكم يتسم بالوعي. إنه يتضمن الإشراف التداولي والتعديل بشأن الآلية الأولى، بما في ذلك التخطيط وتصحيح الأخطاء. وفي منظور شاليس، فإن جدولة التنافس يجري نمذجتها باستخدام المعالجة الموزعة المتوازية، والتحكم التنفيذي يجري نمذجته باستخدام الذكاء الاصطناعي الرمزي.

يمكن رفع مستوى التنشيط لمخططات العمل عن طريق المدخلات الإدراكية. على سبيل المثال، نظرة غافلة (التعرف على الأنماط) من شخص إلى الفراش فور الوصول إلى غرفة النوم يمكن أن تحفز مخطط العمل للخلود إلى النوم، على الرغم من أن النية الأصلية (الخطة) كانت تغيير الملابس.

انطلقت نظرية شاليس عن العمل باستخدام أفكار من الذكاء الاصطناعي (لا سيما نماذج التخطيط)، وقد حظيت بصدى يتفق مع خبرته الإكلينيكية. وقد دعمتها لاحقاً أدلة من فحص للدماغ. كذلك اكتشف علم الأعصاب مؤخرًا عوامل أخرى، ومنها الناقلات العصبية، التي لها دخل في فعل الإنسان. وتلك العوامل يجري تمثيلها الآن في نماذج الكمبيوتر الحالية بناءً على تلك النظرية.

التفاعلات بين جدولة المنافسة والتحكم التنفيذي له صلة بعلم الروبوتات. العامل الذي يتبع خطة ينبغي أن يكون قادرًا على إيقافها أو تغييرها بناءً على ما يلاحظه في

البيئة. تلك الاستراتيجية تميز الروبوتات التي تجمع بين المعالجة «الكائنة» و«التداولية» (انظر الفصل الخامس).

ينبغي لأي مهتم بالذكاء الاصطناعي العام أن يلاحظ أن هؤلاء القلائل من علماء الذكاء الاصطناعي الذين فكروا بجدية في البنية الحاسوبية للعقل ككل يقبلون الأنظمة المختلطة من دون تحفظ. من هؤلاء المهتمين ألين نيويل وأندرسون (ونموذجهما «التوجيه نحو النجاح وتحقيق الأهداف» و«التحكم المتكيف مع التفكير» * اللذان تناولناهما في الفصل الثاني)، وستان فرانكلين (ونموذج «تعلّم وكيل التوزيع الذكي» الخاص بالوعي الموضح في الفصل السادس)، ومينسكي (ونظريته «الاجتماعية» عن العقل)، وآرون سلومان (ونموذج عن محاكاة القلق الذي تناولناه في الفصل الثالث).

باختصار، الأجهزة الافتراضية المنفذة في أدمغتنا تجمع بين نموذج التسلسل والتوازي. والذكاء البشري يتطلب تعاونًا دقيقًا بينهما. وسيفعل الذكاء الاصطناعي العام على المستوى البشري ذلك، إن تحقّق في أي وقت.

الفصل الخامس

الروبوتات والحياة الاصطناعية

تحاكي الحياة الاصطناعية الأنظمة البيولوجية. ومثل الذكاء الاصطناعي بوجه عام، فإن لها أهدافاً تكنولوجية وعلمية. الحياة الاصطناعية مكمل للذكاء الاصطناعي؛ لأن كل أنواع الذكاء التي نعرفها توجد في الكائنات الحية. في الحقيقة، يعتقد كثيرون أن العقل لا ينشأ إلا من الحياة (انظر الفصل السادس). لا يقلق اختصاصيو التكنولوجيا الواقعيون بشأن ذلك السؤال. ولكنهم بالفعل يلتفتون إلى علم الأحياء عندما يريدون ابتكار تطبيقات عملية ذات أنواع متعددة. ومن تلك التطبيقات الروبوتات والبرمجة التطورية والأجهزة الذاتية التنظيم. الروبوتات أمثلة نموذجية على الذكاء الاصطناعي؛ فهي تتمتع بمكانة بارزة، وتمتاز بذكاء ألمعي، كما أنها مهمة للغاية على المستوى التجاري. وعلى الرغم من الاستخدام الواسع النطاق للذكاء الاصطناعي التطوري، فإنه ليس معروفاً بقدر ما تُعرف الروبوتات. كذلك الآلية الذاتية التنظيم أقل شهرةً (باستثناء التعلم غير الموجه؛ انظر الفصل الرابع). ولكن سعيًا لفهم التنظيم الذاتي، فقد أفاد الذكاء الاصطناعي علم الأحياء كما استفاد منه.

الروبوتات الكائنة والحشرات الشائكة

بُنيت الروبوتات منذ قرون على يد ليوناردو دافنشي وغيره. ظهرت إصدارات الذكاء الاصطناعي في خمسينيات القرن العشرين. أذهلت «سلفاة» ويليام جراي وولتر فيما بعد الحرب المراقبين وهي تتفادى العقبات وتبحث عن الضوء. ومن بين أهم أهداف مختبر الذكاء الاصطناعي الذي أقامه معهد ماساتشوستس للتكنولوجيا حديثاً هو بناء «روبوت معهد ماساتشوستس للتكنولوجيا»، بحيث يجمع بين الرؤية الحاسوبية والتخطيط واللغة والتحكم الحركي.

تحقق تقدّم كبير منذ ذلك الحين. والآن، بإمكان الروبوتات تسلق التلال أو السلالم أو الجدران، وبعضها بإمكانه الجري بسرعة أو القفز عاليًا، وبعضها بإمكانه حمل الأحمال الثقيلة ورميها. روبوتات أخرى بإمكانها تفكيك نفسها وتجميع أجزائها مرة أخرى، وأحيانًا تتحول إلى شكل جديد، كأن تتحول إلى دودة (قادرة على اجتياز أنبوب ضيق) أو كرة أو كائن متعدد السيقان (يمشي على أرض مستوية أو وعرة). ما حفز هذا التقدم هو التحول من علم النفس إلى علم الأحياء.

حاكت روبوتات الذكاء الاصطناعي الكلاسيكية أفعال البشر التطوعية. واعتمادًا على نظريات النمذجة العقلية، فقد استخدمت تمثيلات داخلية للعالم وأفعال العامل. ولكنها لم تكن رائعة. ولأنها اعتمدت على التخطيط المجرد، فقد تعرّضت إلى مشكلة الإطار (انظر الفصل الثاني). لم تتفاعل على النحو المطلوب؛ لأن أي تغييرات بيئية حتى لو طفيفة تتطلب التخطيط الاستباقي لإعادة التشغيل، كما أنها لم تتكيف على الظروف الجديدة (التي ليس لها نماذج). كانت الحركة الثابتة صعبة حتى على الأرض المستوية غير الوعرة (ومن هنا جاء اسم روبوت شاكى SHAKEY الذي ابتكره معهد ستانفورد للأبحاث)، ولم تستطع الروبوتات التي سقطت النهوض مرة أخرى. لم تكن ثمة فائدة في كثير من الأبنية، ناهيك عن كوكب المريخ.

الروبوتات في الوقت الحالي مختلفة تمامًا. تحول التركيز من الإنسان إلى الحشرات. ربما لا تكون الحشرات ذكية بالقدر الكافي لنمذجة العالم أو للتخطيط. ولكنها تقدر الآن على ذلك. فسلوكها — السلوك وليس الفعل — ملائم وتكيفي. لكنه في الأساس انعكاسي وليس متعمدًا. إنها تستجيب للموقف الحالي من دون تفكير، ولا تستجيب لاحتمال تخيلي أو حالة مستهدفة. ومن هنا جاءت التسميات: الروبوتات «الكائنة» أو «القائمة على السلوك». (السلوك الكائن ليس محصورًا على الحشرات؛ حدّد علماء النفس الاجتماعي العديد من السلوكيات المرتبطة بالمواقف لدى البشر).

في محاولات إعطاء ردود فعل منعكسة يمكن مقارنتها بأجهزة الذكاء الاصطناعي، يفضل علماء الروبوتات الهندسة على البرمجة. وإذا أمكن، تتجسد ردود الفعل الحسية الحركية ماديًا في تشريح الروبوت، ولا تتوافر في صورة كود برمجي.

إلى أي مدّى يجب أن يتطابق تشريح الروبوت مع تشريح الكائنات الحية؟ لأغراض تكنولوجية، تُقبل الحيل الهندسية المبتكرة. تحتوي الروبوتات اليوم على العديد من الحيل غير الواقعية. ولكن هل يُحتمل أن تتمتع الآليات البيولوجية بكفاءة خاصة؟ لا شك أنها تتمتع بما يكفي من الكفاءات. ولذا يدرس علماء الروبوتات الحيوانات الحقيقية؛ ما الذي

يمكنها فعله (بما في ذلك الاستراتيجيات الملاحية المختلفة)، وما الإشارات الاستشعارية والحركات المحددة المتضمنة، وما الآليات العصبية المسؤولة عن ذلك. يوظف علماء الأحياء بدورهم نمذجة الذكاء الاصطناعي لدراسة هذه الآليات؛ وهو مجال بحثي يُسمى «علم الأعصاب الحاسوبي».

من الأمثلة على ذلك، روبوتات الصرصور الذي طوّرها راندال بير. تمتلك الصراصير ست سيقان متعددة الأجزاء، وهو ما يشير إلى وجود ميزات وعيوب. الحركة على ستة أرجل مستقرة أكثر من الحركة على رجلين (ومفيدة بوجه عام أكثر من العجلات). ولكن يبدو أن التنسيق بين ستة أطراف أصعب من التنسيق بين طرفين. بالإضافة إلى تقرير الساق التي ينبغي أن تتحرك تاليًا، فلا بد للمخلوق أن يجد المكان الصحيح والقوة الكافية والوقت المناسب. وكيف تتفاعل السيقان؟ لا بد أن تكون السيقان مستقلة إلى حد كبير في حركتها؛ لأنه قد لا توجد حصة إلا عند ساق واحدة. ولكن إذا ارتفعت تلك الساق أعلى، فلا بد للسيقان الأخرى أن تكافئ ذلك الارتفاع للحفاظ على التوازن.

تعكس روبوتات بير التشريح العصبي وأدوات التحكم الحسية للصراصير الحقيقية. إنها تستطيع تسلك السلالم والسير على الأرض الوعرة وتجاوز العقبات (بدلاً من مجرد التفادي) والنهوض من بعد السقوط.

تنظر عالمة الروبوتات باربرا ويب إلى صرصور الليل، وليس إلى صرصور المنازل. لم يكن تركيزها على الحركة (ومن ثمّ تستطيع الروبوتات التي طوّرتها استخدام العجلات). بل أرادت من أجهزتها أن تتعرف على أنماط أصوات محددة، وتحدد موقعها وتقرب منها. لا ريب أن هذا السلوك (الاستجابة للصوت) قد ينطوي على العديد من التطبيقات العملية. بإمكان أنثى صرصور الليل أن تفعل ذلك عند سماع أغنية الذكر من فصيلتها. ومع ذلك، لا يمكن لصرصور الليل أن يتعرف إلا على أغنية واحدة لا تغنى إلا بسرعة واحدة وتردّد واحد. تتفاوت السرعة والتردد حسب أنواع صراصير الليل المختلفة. ولكن الأنثى لا تختار بين الأغاني المختلفة، فهي لا تمتلك أدوات استكشاف تشفر مجموعة من الأصوات. إنها تستخدم آلية لا تكون حساسة إلا تجاه تردّد واحد. إنها ليست آلية عصبية، مثل أدوات الاستكشاف السمعية في دماغ الإنسان. إنها أنبوب ذو طول ثابت في الصدر ومتصل بالأذنين في سيقانها الأمامية وبفتحات التنفس. طول الأنبوب هو نسبة دقيقة من الطول الموجي لأغنية الذكر. تؤكد الفيزياء أن عمليات إزاحة الطور (بين الهواء في الأنبوب والهواء الخارجي) لا تحدث إلا في الأغنيات ذات التردد الصحيح، ويعتمد هذا

الفرق في الكثافة اعتمادًا كاملاً على اتجاه مصدر الصوت. أنثى الحشرة مرتبطة عصبياً بالتحرك في ذلك الاتجاه؛ فالذكر يغني وهي تتحرك نحوه. إنه السلوك الكائن بالفعل. اخترت ويب خاصية الاستجابة للصوت لدى صرصور الليل؛ لأن علماء الأعصاب درسوها من كتب. لكن كانت هناك العديد من الأسئلة التي لم يجب عليها؛ هل تتم معالجة اتجاه الأغنية والصوت كلٌّ على حدة أم لا (وكيف)، وهل التعرف على الصوت مستقل عن تحديد مكانه، وكيف تتحفز الحركة لدى الأنثى، وكيف يُتحكم في اتجاهها المتعرج. اخترت ويب أبسط آلية ممكنة (أربع خلايا عصبية فقط) يمكن أن تولّد سلوكاً مشابهاً. بعد ذلك، استخدم نموذجها المزيد من الخلايا العصبية (بناءً على البيانات التفصيلية من الحياة الواقعية)، وتضمنت المزيد من الخصائص العصبية (مثل زمن الاستجابة ومعدل الإطلاق والجهد الغشائي)، وتكامل السماع مع الرؤية. أوضح عملها العديد من المسائل العلمية العصبية، وقد أجاب على بعضها وأثار أسئلة أخرى. ومن ثم أفادت علم الأحياء كما أفادت علم الروبوتات.

(على الرغم من أن الروبوتات أشياء مادية، فإن كثيراً من أبحاث الروبوتات أُجريت بطريقة المحاكاة. على سبيل المثال، روبوتات بير يتم تطويرها «برمجياً» قبل بنائها. وبالمثل، تُصمّم روبوتات ويب في صورة «برامج» قبل اختبارها في العالم الواقعي).

على الرغم من التحول إلى الحشرات في مجال الروبوتات السائد، لم تتوقف الأبحاث بشأن الروبوتات على شكل الإنسان. بعض تلك الروبوتات مجرد دُمى. وبعضها الآخر عبارة عن روبوتات «اجتماعية» — أو مرافقة — مصممة كي يستخدمها كبار السن والمعاقين في تنفيذ المهمات المنزلية (انظر الفصل الثالث). ليس الهدف من تصميم هذه الروبوتات أن يكونوا عبيداً يؤديون مهام ثانوية، بل أن يكونوا مساعدين شخصيين مستقلين. بعضها له مظهر «جميل» ورموش طويلة وأصوات جذابة. يمكنها أن تتواصل بالعين مع المستخدمين، وتتعرف على وجوه الأشخاص وأصواتهم. كذلك يمكنها — إلى حد ما — أن تُجري محادثات غير مكتوبة، وتفسّر الحالة العاطفية للمستخدم، وتولد ردود فعل «عاطفية» (بحيث تظهر تعبيرات وجه أو أنماط كلام مشابهة للبشر) بأنفسها.

على الرغم من أن بعض الروبوتات كبيرة الحجم (لحمل الحمولات الثقيلة أو للمشبي على الأرض الوعرة)، فإن معظمها صغير الحجم. وبعضها الآخر بالغ الصغر، حيث تُستخدم داخل الأوعية الدموية على سبيل المثال. وكثيراً ما ترسل للعمل بأعداد كبيرة. وعندما تشترك عدة روبوتات في مهمة، تثار الأسئلة بشأن طريقة تواصلها (إذا أمكن)، وكيف يمكن ذلك المجموعة من إنجاز مهام لا يمكن أن ينجزها فرد واحد.

للإجابة على هذه الأسئلة، غالبًا ما يفكر علماء الروبوتات في الحشرات الاجتماعية مثل النمل والنحل. تلك الأنواع تُضَرَّبُ مثالاً على «المعرفة الموزعة» (انظر الفصل الثاني)، وتلك الخاصية تتوزع فيها المعرفة (والعمل الملائم) على المجموعة بكاملها بدلاً من أن تتوافر لحيوان واحد.

إذا كانت الروبوتات بالغة البساطة، فقد يتحدث مطوِّروها عن «ذكاء السرب»، ويحللون أنظمة الروبوتات التعاونية باعتبارها إنساناً آلياً خلوياً. الإنسان الآلي الخلوي عبارة عن نظام لوحات فردية، وكل وحدة تأخذ حالة من عدد حالات باتباع قواعد بسيطة؛ إذ تعتمد على الحالة الحالية للوحدات المجاورة. قد يكون النمط الكلي لسلوك الإنسان الآلي الخلوي معقداً إلى حد الذهول. التماثل الأساسي هو الخلايا الحية المتعاونة لدى الكائنات المتعددة الخلايا. تتضمن العديد من إصدارات الذكاء الاصطناعي الخوارزميات المتدفقة التي تُستخدم في حشود الخفافيش أو الديناصورات في أفلام الرسوم المتحركة بهوليود.

ينطبق مفهوم المعرفة الموزعة ومفهوم ذكاء السرب على الإنسان أيضاً. يُستخدم مفهوم ذكاء السرب عندما لا تكون المعرفة المعنية شيئاً يمكن أن يحوزه فرد واحد من المجموعة (مثل السلوك العام للحشود الضخمة). وغالبًا ما يُستخدم مصطلح المعرفة الموزعة عندما يمكن للأفراد المشاركين أن يحوزوا كل المعرفة ذات الصلة، ولكن لا يتفرد بها أحد. على سبيل المثال، أوضح عالم الأنثروبولوجيا إدوين هاتشين كيف أن معرفة الملاحة يشاركها أفراد طاقم السفينة، كما أنها تتجسد في أشياء مادية مثل الخرائط (وموقع) جداول الرسوم البيانية.

الحديث عن تجسيد المعرفة في أشياء مادية قد يبدو غريباً أو مجازياً في أفضل الحالات. ولكن هناك الكثير اليوم ممن يزعمون أن عقول البشر تتجسد حرفياً، ليس فقط في أفعال الناس الجسدية، ولكن أيضاً في المصنوعات الثقافية التي يتعاملون معها في العالم الخارجي. تأسست نظرية «العقل المتجسد/الخارجي» جزئياً في عمل أنجزه رائد تحول الروبوتات من شكل الإنسان إلى الحشرات، ألا وهو رودني بروكس من معهد ماساتشوستس للتكنولوجيا.

الآن، بروكس من مطوِّري الروبوتات الأساسيين في الجيش الأمريكي. وفي ثمانينيات القرن العشرين، كان لا يزال عالم روبوتات ناشئاً، ولكن أُحبط من عدم جدوى مخططي نمذجة العالم في الذكاء الاصطناعي الرمزي. تحول إلى الروبوتات الكائنة لأسباب

تكنولوجية بحتة، ولكن سرعان ما حول نهجه إلى نظرية عن السلوك التكيفي بوجه عام. تلك النظرية ذهبت إلى ما هو أبعد من الحشرات؛ وقد احتج قائلاً إنه حتى أعمال الإنسان لا تتضمن تمثيلات داخلية. (أو يقول ضمناً إنها لا تنطوي على تمثيلات عادةً).
نقده للذكاء الاصطناعي الرمزي أعجب علماء النفس والفلاسفة. بعضهم أبدى تأييداً كبيراً. سبق أن أشار علماء النفس إلى أنه كثيراً ما يكون سلوك الإنسان مرهوناً بالموقف، كأن يكون له دور في بيئات اجتماعية مميزة. كذلك ألقى علماء النفس المعرفي الضوء على الرؤية المتحركة، وفيها تكون الحركة الجسدية للعامل هي المفتاح. واليوم، أصبح لنظريات العقل المجسد تأثير كبير في غير الذكاء الاصطناعي (انظر الفصل السادس).
لكن البعض الآخر أمثال ديفيد كيرش كانوا — ولا يزالون — يعارضون بشدة، بحجة أن التمثيلات التركيبية ضرورية في أنواع السلوكيات التي تتضمن مفاهيم. ومن الأمثلة على ذلك التعرف على عدم التباين الإدراكي، وفيه يمكن التعرف على كائن من عدة مشاهد مختلفة، وإعادة تعريف الأفراد بمرور الوقت، وضبط النفس الاستباقي (التخطيط)، والتفاوض بشأن الأهداف المتضاربة وليس مجرد جدولتها، والتفكير المنطقي في الواقع المضاد، واللغة. تعترف هذه الانتقادات أن الروبوتات الكائنة تُظهر سلوكاً خالياً من المفاهيم لدرجة أنها انتشرت بدرجة أكبر مما اعتقد العديد من الفلاسفة. وعلى الرغم من ذلك، يتطلب كل من المنطق واللغة والتصرفات البشرية المدروسة الحوسبة الرمزية.
يرفض العديد من علماء الروبوتات أيضاً ادعاءات بروكس الأكثر تطرفاً. مجموعة آلان ماكورث واحدة من المجموعات العاملة في كرة القدم الروبوتية، وتشير إلى «المداولات التفاعلية» التي تتضمن الإدراك الحسي واتخاذ القرارات في الوقت الفعلي والتخطيط والتعرف على الخطة والتعلم والتنسيق. إنها تسعى إلى تكامل الذكاء الاصطناعي التقليدي الجميل ووجهات النظر الكائنة. (ولذلك، تبني المجموعة أنظمة مختلطة؛ انظر الفصل الرابع).

بوجه عام، التمثيلات بالغة الأهمية في عملية اختيار الإجراءات لدى الروبوتات، ولكن ليست بتلك الأهمية فيما يتعلق بتنفيذ الإجراءات. لذا، فإن المهرجين الذين قالوا إن «الذكاء الاصطناعي» يشير إلى «الحشرات الاصطناعية» لم يوفقوا في ذلك تماماً.

الذكاء الاصطناعي التطوري

يعتقد معظم الناس أن الذكاء الاصطناعي يستلزم تصميمًا يراعي أدق التفاصيل. وبالنظر إلى طبيعة الكمبيوتر التي لا ترحم، فكيف يكون الأمر غير ذلك؟ حسنًا، قد يكون.

على سبيل المثال، نتجت الروبوتات التطورية (التي تتضمن قدرًا من الروبوتات الكائنة) من الجمع بين الهندسة/ البرمجة الدقيقة والتباين العشوائي. تتلخص المسألة في تطورها غير المتوقع، وليس في العناية بتصميمها.

بوجه عام، يمتاز الذكاء الاصطناعي التطوري بهذه السمة. لقد بدأ في الذكاء الاصطناعي الرمزي، ولكنه يُستخدم أيضًا في الترابطية. تطبيقاته العملية المتعددة تتضمن الفنون (يمكن أن تنطوي على عدم القدرة على التنبؤ) وتطوير أنظمة حساسة تجاه السلامة مثل محركات الطائرات.

يمكن أن يغير البرنامج نفسه (بدلاً من أن يعيد المبرمج كتابته)، ويمكن أن يحسن من نفسه حتى عن طريق استخدام الخوارزميات الوراثية. إنها مستوحاة من علم الوراثة الواقعي، وإنها تمكّن كلاً من التباين العشوائي والانتقاء غير العشوائي. الانتقاء يتطلب معيار النجاح، أو «دالة الصلاحية» (المشابهة للانتقاء الطبيعي في علم الأحياء) التي تعمل جنباً إلى جنب مع الخوارزميات الوراثية. تحديد دالة الصلاحية أمر بالغ الأهمية.

في البرامج التطورية، لا يستطيع البرنامج الأولي الموجه حسب المهمة حل المهمة بكفاءة. قد لا يستطيع حلها على الإطلاق؛ فقد تكون مجموعة من التعليمات غير المتسقة أو شبكة عصبية ذات اتصالات عشوائية. لكن البرنامج الشامل يتضمن الخوارزميات الوراثية في الخلفية. وهذه الخوارزميات تغير القواعد الموجهة بالمهام. تتشابه التغيرات التي تحدث عشوائياً مع الطفرة النقطية والعبور في علم الأحياء. لذا قد يتغير رمز واحد في تعليمة مبرمجة، أو قد «تتبادل» سلسلة رموز قصيرة بين تعليمتين.

تعقد المقارنة بين برامج المهمات المتعددة ضمن أي جيل، وتُستخدم البرامج الأنجح لتوليد الجيل التالي. كذلك يمكن الاحتفاظ ببضعة برامج أخرى (تُختار عشوائياً) بحيث لا تُفقد الطفرات المفيدة التي لم يكن لها تأثير جيد حتى الآن. بمرور الأجيال، تزيد كفاءة برنامج المهمة. ويوجد الحل الأمثل في بعض الأحيان. (في بعض الأنظمة التطورية، تُحل مسألة إحالة الاستحقاق (انظر الفصل الرابع) بإدخال بعض التباينات في خوارزمية لواء الدلو، حيث إن تلك الخوارزمية لا تحدد إلا الأجزاء المسؤولة عن النجاح في البرنامج التطوري المعقد).

بعض أنظمة الذكاء الاصطناعي التطوري آلية بالكامل؛ بمعنى أن البرنامج يطبق دالة الصلاحية في كل جيل، ويترك كي يتطور من دون أن يشرف عليه أحد. وهنا، يجب أن تكون المهمة بالغة الوضوح؛ باستخدام فيزياء محركات الطائرة على سبيل المثال. وعلى

الجانب الآخر، عادةً ما يكون الفن التطوري ذا درجة تفاعل عالية (يختار الفنان الأفضل في كل جيل)؛ لأنه لا يمكن ذكر دالة الصلاحية — المعايير الفنية بوضوح — ذكرًا واضحًا. تتمتع معظم الروبوتات التطورية بالقدرة على التفاعل في أوقات متقطعة. يتطور تشريح الروبوت (مثل أجهزة الاستشعار والتوصيلات الحسية الحركية) ووحدة التحكم (الدماغ) أو أحدهما تلقائيًا، ولكن في مرحلة المحاكاة. وفي معظم الأجيال، لا يوجد روبوت مادي. لكن عند كل ٥٠٠ جيل، ربما يخضع التصميم المتطور لاختبار في جهاز مادي. عادةً لا تحيا الطفرات التي لا فائدة منها. اكتشف فريق إنمان هارفي بجامعة ساسكس (عام ١٩٩٣) أن إحدى عيني الروبوت وكل «شعيرات الاستشعار» قد تفقد اتصالاتها الأولية بالشبكة العصبية المتحركة، إذا لم تكن المهمة بحاجة إلى أي من الرؤية المتعمقة ولا اللمس. (وبالمثل، في الحوسبة البصرية، تُستخدم القشرة السمعية في الأشخاص الذين يعانون صممًا وراثيًا أو في الحيوانات المحرومة من حاسة السمع؛ بمعنى أن الدماغ يتطور على مدار الحياة، وليس عبر الأجيال فحسب).

قد ينطوي الذكاء الاصطناعي التطوري على مفاجآت عميقة. على سبيل المثال، يخضع الروبوت الكائن إلى تطويرات (بجامعة ساسكس أيضًا) لتوليد حركة تتحاشى العقبات تجاه هدف، وقد طور أداة لاستكشاف الاتجاهات مشابهة للتي توجد في الدماغ. تضمن عالم الروبوت مثلثًا أبيض من الورق المقوّى. وعلى نحو غير متوقّع، نشأت شبكة صغيرة متصلة عشوائيًا في وحدة التحكم، وقد استجابت إلى تدرج ضوء/ظلام في اتجاه معين (جانب من المثلث). عندئذٍ، تتطور تلك الشبكة وتصبح جزءًا لا يتجزأ من آلية بصرية حركية واتصالاتها بوحدات الحركة (العشوائية في البداية)؛ مما يمكّن الروبوت من استخدام الكائن باعتباره أداة مساعدة على الملاحظة. لا تعمل الآلية مع المثلثات ذات اللون الأسود، ولا في الجانب المعاكس. كانت الآلية قائمة بذاتها، فلا يوجد نظام شامل لأدوات استكشاف التوجيه. ولكنها كانت مفيدة على أي حال. تكرّرت تلك النتيجة المحيّرة كثيرًا. وباستخدام الشبكات العصبية من الأنواع المختلفة، وجد فريق جامعة ساسكس أن كل حل ناجح قد طور أداة نشطة لاستكشاف الاتجاهات؛ لذلك لم تتغير الاستراتيجية السلوكية العالية المستوى. (تباينت تفاصيل التنفيذ الدقيقة، لكنها غالبًا ما كانت متشابهة جدًا).

في مناسبة أخرى، استخدم فريق جامعة ساسكس الخوارزميات الوراثية لتصميم الدوائر الكهربائية في الأجهزة. تمثّلت المهمة في تطوير أجهزة ذبذبة. ومع ذلك، ما ظهر كان

مستشعرًا بدائيًا لموجات الراديو، وكان يلتقط الإشارة الأرضية من شاشة كمبيوتر قريبة. وقد اعتمد ذلك على معلومات فيزيائية غير متوقعة. كان بعضها متوقعًا (الخصائص الشبيهة بالهواء لجميع اللوحات الكهربائية المطبوعة)، على الرغم من أن الفريق لم يفكر فيها مسبقًا. ولكن البعض الآخر كان عرضيًا، ويبدو أنها ليست ذات صلة. تضمنت المعلومات القرب المكاني من شاشة الكمبيوتر، والترتيب الذي تُضبط به المفاتيح التماثلية، وحقيقة ترك مكواة اللحام على مقربة من طاولة عمل وهي موصلة بالتيار الكهربائي. (لم تتكرر تلك النتيجة؛ ففي المرة التالية، قد يتأثر هوائي الراديو بكيمياء ورق الحائط).

مستشعر موجات الراديو مثير للإعجاب؛ لأن العديد من علماء الأحياء (والفلاسفة) يقولون إنه لا ينشأ شيء جديد جذريًا من الذكاء الاصطناعي، حيث إن كل نتائج برامج الكمبيوتر (بما في ذلك النتائج العشوائية للخوارزميات الوراثية) يجب أن تقع ضمن حيز الاحتمالات التي يحددها البرنامج. يقولون كذلك إنه لا يمكن لغير التطور البيولوجي أن يولد مستشعرات إدراكية. إنهم يصرحون بأنه يمكن أن يتطور مستشعر بصري ضعيف للذكاء الاصطناعي ويتحول إلى مستشعر أفضل. لكنهم يقولون إن المستشعر البصري الأول لا يمكن أن ينشأ إلا في عالم مادي تحكمه السببية. فالطفرة الجينية العشوائية التي تدفع مادة كيميائية حساسة تجاه الضوء يمكن أن تُدخل الضوء — الموجود بالفعل في العالم الخارجي — إلى بيئة الكائن الحي. ولكن بالمثل، أدخل مستشعر راديو غير متوقع موجات الراديو إلى بيئة الجهاز. وبالفعل يعتمد الأمر جزئيًا على السببية الفيزيائية (المكونات الإضافية وغيرها). لكنه كان تدريبيًا في الذكاء الاصطناعي، وليس في علم الأحياء. تتطلب الحداثة الجذرية تأثيرات خارجية بالفعل؛ لأنه صحيح أن البرنامج لا يمكنه تجاوز مساحة احتماليته. ولكن ليس بالضرورة أن تكون هذه المؤثرات مادية. فنظام الخوارزميات الوراثية المتصل بالإنترنت قد يطور تحديثات جوهرية عن طريق التفاعل مع العالم الافتراضي.

وقعت مفاجأة أخرى قبل ذلك بكثير في مجال الذكاء الاصطناعي التطوري؛ مما شجّع الأبحاث الجارية على التطور بالمعنى الحرفي للكلمة. استخدم عالم الأحياء توماس راي الخوارزميات الوراثية لمحاكاة بيئة الغابات الاستوائية المطيرة. لقد رأى التطور التلقائي للطفيليات، والمقاومة لتلك الطفيليات، والطفيليات الفائقة القدرة على التغلب على تلك المقاومة. اكتشف أيضًا أن «الطفرات» المفاجئة في التطور (الظاهري) يمكن توليدها بتعاقب طفرات (جينية) كامنة. بالطبع يؤمن أتباع داروين الأرثوذكس بهذا.

ولكن من غير البديهي أن يقول بعض علماء الأحياء مثل ستفين جاي جولد أن العمليات غير الداروينية لا بد أنها داخلية في تلك الطفرات هي الأخرى. في الوقت الراهن، تتباين معدلات الطفرات المحاكاة تبايناً منهجياً، ويجري تتبعها، ويحلل باحثو الخوارزميات الوراثية «مشاهد الصلاحية» و«الشبكات العصبية» [منقولة من دون تعديل] و«الانحراف الجيني». هذا العمل يبين كيف يمكن الحفاظ على الطفرات على الرغم من أنها لم تزد صلاحية التكاثر (حتى الآن). ومن ثم يساعد الذكاء علماء الأحياء على وضع نظريات تطويرية بوجه عام.

التنظيم الذاتي

الميزة الأساسية في الكائنات الحية هي قدرتها على بناء نفسها. التنظيم الذاتي هو ظهور النظام تلقائياً من أصل ذي درجة نظام أقل. إنها خاصية محيرة، وتكاد تكون متناقضة. وليس واضحاً إن كان يمكن حدوثها في الجمادات أم لا.

بوجه عام، التنظيم الذاتي ظاهرة إبداعية. تناولنا الإبداع النفسي (كل من التاريخي والفردى) في الفصل الثالث، والتعلم الترابطي (غير الموجّه) الذاتي التنظيم في الفصل الرابع. ونركز في هذا الفصل على أنواع التنظيم الذاتي المدروس في الأحياء. تتضمن الأمثلة تطور السلالات (ضرب من الإبداع التاريخي)، وتكوين الأجنة والتحول (المشابه للإبداع الفردى في علم النفس)، ونمو الدماغ (إبداع فردى يتبعه إبداع تاريخي)، وتكوين الخلايا (إبداع تاريخي عندما تدب الحياة في الأوصال ثم يتحول إلى إبداع فردى). كيف يساعدنا الذكاء الاصطناعي على فهم تلك الأمور؟

شرح آلان تورينج التنظيم الذاتي عام ١٩٥٢ بأنه العودة إلى الأساسيات. تساءل كيف لشيء متجانس (مثل البويضة غير المتميزة) أن يؤسس بنية. اعترف أن قدرًا كبيراً من النمو البيولوجي يضيف تنظيمًا جديدًا إلى النظام الموجود مسبقاً؛ مثل سلسلة التغيرات التي تحدث في القناة العصبية للجنين. ولكن تولد النظام من التجانس هو الحالة الأساسية (وأبسطها رياضياً).

طرح علماء الأجنة بالفعل فكرة «أدوات التنظيم»؛ وهي مواد كيميائية غير معروفة توجه النمو بطرق غير معروفة. كذلك لم يستطع تورينج تحديد أدوات التنظيم تلك. بل إنه درس مبادئ عامة عن الانتشار الكيميائي.

أظهر أنه إذا تقابلت الجزيئات المختلفة، فستعتمد النتيجة على معدلات الانتشار وتركيزاتها والسرعات التي يمكن لتفاعلاتها أن تهدم/تبني جزيئاتها بها. لقد فعل ذلك

عن طريق تغيير الأعداد في معادلات كيميائية تخيلية ودراسة النتائج. لم تنتج بعض مجموعات الأرقام إلا مخاليط من المواد الكيميائية ولكن لا شكل لها. ولكن تولّد نظام من مجموعات أخرى، مثل القمم المنتظمة لتركيز جزيء بعينه. قال إن هذه القمم الكيميائية يمكن التعبير عنها بيولوجياً بأنها علامات سطحية (شرائط)، أو بأنها أصول لبنيات متكررة مثل البتلات أو أجزاء الجسم. يمكن أن تؤدي تفاعلات الانتشار في ثلاثة أبعاد إلى تفريغ، مثل تكون المعيدة عند الجنين.

أُعيد تنظيم هذه الأفكار على الفور باعتبارها مثيرة للاهتمام بدرجة كبيرة. لقد حلّت لغزًا كان مستعصياً فيما سبق بشأن تولد النظام من رحم أصل غير منتظم. في خمسينيات القرن العشرين، لم يستطع علماء الأحياء أن يفعلوا الكثير إزاء هذا اللغز. فقد اعتمد تورينج على التحليل الرياضي. أجرى بالفعل قدرًا من المحاكاة (الشاقة للغاية) يدويًا، وأتبعها بنمذجة على جهاز كمبيوتر بدائي. لكن لم يكن الجهاز به قوة حاسوبية كافية لإجراء المجاميع ذات الصلة، أو لاستعراض تباينات الأعداد بطريقة منهجية. كذلك لم تكن رسومات الكمبيوتر متاحة كي تحول قوائم الأعداد إلى شكل واضح ومفهوم.

اضطرَّ كلُّ من الذكاء الاصطناعي وعلم الأحياء أن ينتظرا ٤٠ سنة قبل أن تتطور رؤى تورينج. استعرض خبير رسومات الكمبيوتر جريج تورك معادلات تورينج، وكان في بعض الأحيان «يجمد» نتائج المعادلة قبل أن يطبق معادلة أخرى. يذكّرنا هذا الإجراء بتشغيل/إيقاف تشغيل التبديل بين الجينات، ويضرب مثالاً على تولد النمط من النمط؛ وهو ما ذكره تورينج ولكن لم يستطع تحليله. في نموذج الذكاء الاصطناعي الذي طوره تورك، لم تولّد معادلات تورينج شرائط وعلامات الدلطيقي (كما فعلت عمليات المحاكاة اليدوية)، بل ولدت بقع النمر والفهد وشبكيات الزرافة وأنماط سمكة التنين.

استخدم باحثون آخرون سلاسل معادلات أعقد؛ ومن ثم أدت إلى أنماط أعقد. بعضهم كانوا علماء في الأحياء النمائي ممن يعرفون المزيد عن الكيمياء الحيوية الفعلية. على سبيل المثال، درس براين جودوين دورة حياة طحالب أسيتابولاريا. هذا الكائن الأحادي الخلية يحوّل نفسه من نقطة لزجة لا شكل لها إلى ساق ممدود، وبعد ذلك تنمو قمة مسطحة، ثم تنمو حلقة من العقد حول الحافة، وبعدها تنمو هذه العقد إلى نتوءات جانبية وفروع، وفي النهاية تتحد النتوءات الجانبية وتكوّن غطاءً يشبه المظلة. أوضح خبراء الكيمياء الحيوية أنه يدخل في تلك العملية ما يزيد على ٣٠ معلمة أيضاً (مثل تركيزات الكالسيوم، والتقارب بين الكالسيوم وبروتينات معينة، والمقاومة الميكانيكية

للهيكل الخلوي). أجرى نموذج أسيتابولاريا الحاسوبي الذي طوره جودوين حلقات تغذية راجعة معقدة ومتكررة، حيث يمكن لتلك المعلومات أن تتغير ما بين لحظة وأخرى. نتج عن ذلك تحولات جسدية عديدة.

وعلى غرار تورينج وتورك، تلاعب جودوين بالقيم العددية كي يرى أي القيم التي ستولد أشكالاً جديدة بالفعل. إنه لم يستخدم سوى الأعداد ضمن النطاقات المرصودة في الكائن الحي، ولكنها كانت عشوائية.

قد وجد أن أنماطاً معينة يتكرر نشوءها، مثل التبديل بين تركيزات الكالسيوم العالية/المنخفضة عند حافة الساق (التمائل الناشئ من الدوارة). لم تعتمد تلك الأنماط على خيار محدد من قيم المعلومات، ولكنها كانت تظهر تلقائياً إذا عُينت القيم في أي مكان ضمن نطاق كبير. إضافة إلى ذلك، بمجرد أن تبزغ الدوارة، فإنها تستمر. وعلى حد قول جودوين، فقد تصبح أساساً لتحويلات تؤدي إلى ميزات متكررة الحدوث. قد يحدث هذا في تطور السلالات كما يحدث في التطور الفردي (إبداع تاريخي وكذلك إبداع فردي)، ومن الأمثلة على ذلك تطور ذوات الأربع.

لم ينشأ مطلقاً غطاء مظلة من هذا النموذج. من المحتمل أن يتطلب هذا معلومات إضافية، بحيث تمثل تفاعلات كيميائية غير معروفة حتى الآن داخل طحالب أسيتابولاريا الحقيقية. أو ربما تقع تلك الأغشية ضمن مساحة الاحتمالية للنموذج؛ ومن ثمّ يمكن أن تنشأ منه من حيث المبدأ، ولكن فقط إذا كانت القيم العددية محدودة للغاية لدرجة أنه من غير المحتمل العثور عليها عن طريق البحث العشوائي. (لم تنشأ النتوءات الجانبية أيضاً، ولكن السبب في ذلك ببساطة راجع إلى ضعف القوة الحاسوبية؛ أي إن البرنامج بأكمله سيحتاج إلى أن ينفَّذ على مستوى أقل لكل نتوء جانبي فردي).

رسم جودوين مغزى نظرياً مثيراً للاهتمام. رأى الدورات على أنها أشكال «شاملة» تحدث في العديد من الحيوانات والنباتات، وذلك على خلاف المظلة. هذا يقترح أن تلك الدورات ليست ناتجة عن آليات كيميائية حيوية بالغة التحديد توجّهها جينات تتطور بشكل مفاجئ، ولكنها ناتجة عن العمليات العامة (مثل انتشار التفاعل) الموجودة لدى معظم الكائنات الحية أو كلها. ربما تشكل هذه العمليات أساس علم الأحياء «البنائي»؛ وهو علم عام من علوم الأشكال، وتفسيراته ربما تسبق نظرية الانتقاء التي وضعها داروين، على الرغم من أنها تتسق معها تماماً. (كان الاحتمال ضمناً في مناقشة تورينج، وأكّد عليها دارسو طومسون، وهو عالم أحياء استشهد بذلك الاحتمال، ولكن تورينج نفسه تجاهله).

يعمل انتشار التفاعل حسب القوانين الفيزيائية الكيميائية التي تحدد التفاعلات الموضوعية بين الجزيئات؛ بمعنى أنه يمكن تمثيل تلك القوانين في الإنسان الآلي الخلوي. عندما حدّد جون فون نيومان الإنسان الآلي الخلوي، أشار إلى أنه يمكن تطبيقه على الفيزياء من حيث المبدأ. وفي الوقت الراهن، يستخدم باحثو الحياة الاصطناعية الإنسان الآلي الخلوي في عدة أغراض، لا سيما جيل الأنماط البيولوجية ذات الصلة. على سبيل المثال، الإنسان الآلي الخلوي البالغ التبسيط (المحدد على بُعد واحد فقط) بإمكانه توليد أنماط مشابهة للحياة إلى حد كبير، مثل الأنماط في الصدف.

ربما يكون من المثير للاهتمام على وجه الخصوص استخدام الحياة الاصطناعية للإنسان الآلي في محاولة لوصف «الحياة كما يمكن أن تكون»، وليس فقط «الحياة كما نعرفها». استعرض كريستوفر لانجتون (الذي أطلق اسم «الحياة الاصطناعية» عام ١٩٨٧) العديد من الأناس الآليين المحددين عشوائياً، وأخذ يلاحظ ميلها إلى توليد النظام. ولكن العديد منها لم ينتج غير الفوضى. وبعضها الآخر شكّل بنيات متكررة مملّة أو حتى ثابتة. ولكن قلة منها ولّد أنماطاً لا تنفك عن التغيير وإن كانت ثابتة نسبياً، وهذا ما يميز الكائنات الحية على حد قول لانجتون (ويميز الحوسبة أيضاً). وما يثير الدهشة أن هؤلاء الأناس الآليين يتشاركون القيم العددية نفسها على قياس بسيط من التعقيد المعلوماتي للنظام. اقترح لانجتون أن «معلمة لامدا» تلك ربما تنطبق على كل الكائنات الحية، سواء على الأرض أو المريخ.

لا يشكّل التنظيم الذاتي الأجسام الكاملة فحسب، بل أعضاء الجسم أيضاً. الدماغ على سبيل المثال يتطور بفضل عمليات تطورية (على مدار العمر وعبر الأجيال) كما يتطور بالتعلم غير الموجه. ربما يكون لذلك التعليم نتائج تمييزية إلى حد بعيد (إبداع تاريخي). لكن النمو المبكر للدماغ لدى كل فرد يخلق بنيات عصبية يمكن التنبؤ بها. على سبيل المثال، تمتلك القروود الحديثة الولادة أدوات استكشاف للاتجاهات تدور في ٣٦٠ درجة دوراناً منهجياً. وما كانت لتتعلم ذلك من التجربة في العالم الخارجي؛ لذا من الطبيعي أن يُفترض أنها مشفرة في الجينات. لكن الأمر ليس كذلك. بل إنها تنشأ تلقائياً من شبكة عشوائية في البداية.

لم تظهر تلك الخاصية بالنمذجة الواقعية الأحيائية باستخدام الكمبيوتر على يد علماء الأعصاب، بل أيضاً باستخدام الذكاء الاصطناعي «الخالص». عرّف الباحث رالف لينسكي من شركة آي بي إم شبكات التغذية الأمامية المتعددة الطبقات (انظر الفصل الرابع)،

وأوضح أن قواعد هيب البسيطة — في ضوء عشوائية النشاط — (مثل الضوضاء داخل دماغ الجنين) يمكن أن تولّد مجموعات منظمة من أدوات استكشاف الاتجاهات. لم يعتمد لينسكر على الشروحات العملية وحدها، ولم يركز على أدوات استكشاف الاتجاهات وحدها؛ وتنطبق نظرية «إنفوماكس» التجريدية التي وضعها على أي شبكة من هذا النوع. إنها تنص على أن اتصالات الشبكة تتطور لتعظيم كمية المعلومات المحفوظة عند تحويل الإشارات في كل مرحلة معالجة. تتشكل كل الاتصالات في ظل قيود تجريبية معينة، مثل القيود البيوكيميائية والتشريحية. ولكن الرياضيات تضمن أن يظهر نظام تعاوني من الوحدات المتواصلة. ترتبط نظرية إنفوماكس بتطور السلالات أيضًا. إنه يجعل من غير البديهي أن يكون هناك طفرة واحدة تكيفية في ظل تطور نظام معقد. تتبخر الحاجة الواضحة للعديد من الطفرات المتزامنة إذا كان كل مستوى يمكن أن يتكيف تلقائيًا مع تغيير طفيف في مستوى آخر.

بالنسبة إلى التنظيم الذاتي على المستوى الخلوي، فقد تم نمذجة كل من الكيمياء الحيوية داخل الخلايا وتكوين الخلايا/جدار الخلايا. وهذا العمل يستغل عمل تورينج في انتشار التفاعل. ولكنه يعتمد على المفاهيم الأحيائية أكثر من اعتماده على الأفكار المترسّخة في الحياة الاصطناعية.

باختصار، يقدم الذكاء الاصطناعي العديد من الأفكار النظرية المتعلقة بالتنظيم الذاتي. وكذلك تكثر الأعمال الفنية الذاتية التنظيم.

الفصل السادس

لكن هل هو ذكاء حقاً؟

لنفترض أن أداء أنظمة الذكاء الاصطناعي العام المستقبلية (المزودة بشاشة أو الروبوتات) تساوى مع أداء البشر. هل ستتمتع تلك الأنظمة بذكاء حقيقي أو فهم حقيقي أو إبداع حقيقي؟ هل سيكون لديها ذوات أو مكانة أخلاقية أو اختيار حر؟ هل سيكون لها وعي؟ ولو من دون وعي، فهل يمكن أن تتمتع بأي من الخصائص الأخرى؟

إنها ليست أسئلة علمية، بل إنها أسئلة فلسفية. يشعر كثيرون بديهياً أن الإجابة على كل سؤال هي «من الواضح، لا!»

لكن الأمور لا تستقيم على الدوام. إننا نريد حججاً دقيقة، وليس مجرد بديهيات مدروسة. ولكن تُظهر هذه الحجج أنه لا يوجد إجابات عن تلك الأسئلة يتعذر دحضها. يرجع السبب إلى أن المفاهيم المتضمنة جدلية إلى حد كبير. فقط إذا فُهمت تلك الأسئلة على نحو مُرضٍ، فيمكن أن نطمئن إلى كون الذكاء الاصطناعي العام الافتراضي — أو عدم كونه — ذكياً. باختصار، لا أحد يعلم على وجه اليقين.

قد يقول البعض هذا لا يهم، ما يمكن أن تفعله أنظمة الذكاء الاصطناعي العام هو المهم. لكن يمكن أن تؤثر الإجابات في مدى ارتباطنا بها كما سنرى.

إذن، لن يقدم هذا الفصل إجابات جلية. ولكنه سيقتراح أن بعض الإجابات منطقية أكثر من غيرها. وسيوضح كيف استخدم (بعض) الفلاسفة مفاهيم الذكاء الاصطناعي لإلقاء الضوء على طبيعة العقول الحقيقية.

اختبار تورينج

في ورقة بحثية نُشرت في مجلة «مايند» الفلسفية عام ١٩٥٠، وصف ألان تورينج ما أسماه اختبار تورينج. يطرح الاختبار سؤالاً عما إن كان بإمكان المرء أن يميز هل يتفاعل مع

جهاز كمبيوتر أم مع شخص بعد ٣٠ بالمائة من مدة تصل إلى خمس دقائق أم لا. قال ضمناً إنه إن لم يتحقق ذلك، فلن يكون هناك سبب لإنكار أن الكمبيوتر بإمكانه التفكير حقاً.

كان هذا استهزاءً. على الرغم من تصدّر اختبار تورينج الصفحات الافتتاحية، فإنه كان عنصرًا ثانويًا في ورقة بحثية كان الهدف الأساسي منها أن تكون بيانًا عن الذكاء الاصطناعي. في الحقيقة، ذكر تورينج أمر هذا الاختبار لصديقه روبن غاندي باعتبارها «دعاية» غير جادة، بل دعوة إلى الضحك وليس النقد الجاد.

وعلى الرغم من ذلك، نهض الفلاسفة إلى النقد. قال معظمهم إنه حتى إن لم تفرق ردود فعل البرنامج عن ردود فعل الإنسان، فهذا لا يثبت ذكاءه. وكان الاعتراض الأشهر — ولا يزال — هو أن اختبار تورينج لا يهتم إلا بالسلوك المرصود؛ ومن ثمّ يمكن أن يجتازه زومبي، وهو شيء يتصرف مثلنا ولكنه يفتقر إلى الوعي.

يفترض هذا الاعتراض أن الذكاء يتطلب الوعي، وأن مخلوقات الزومبي محتملة من حيث المنطق. سنرى أن بعض حسابات الوعي (في قسم «الذكاء الاصطناعي والوعي المدرك بالحواس») تشير إلى أن مفهوم الزومبي ليس مُتماسكًا. وإذا كانوا على حق، فإنه لا يوجد ذكاء اصطناعي عام يمكن أن يكون زومبي. وفي هذا الشأن، سيكون لاختبار تورينج مسوغاته.

أثار اختبار تورينج اهتمام الفلاسفة (والعامة) كثيرًا. لكنه لم يكن مهمًا في الذكاء الاصطناعي. غالبًا ما يهدف الذكاء الاصطناعي إلى تقديم أدوات مفيدة، ولا يهدف إلى محاكاة ذكاء الإنسان، بل إلى جعل المستخدمين يعتقدون أنهم يتعاملون مع إنسان. والحق أن الباحثين في الذكاء الاصطناعي التّواقيين إلى الشهرة يزعمون في بعض الأحيان ويسمحون للصحافة أن تزعم أن أنظمتهم اجتازت اختبار تورينج. ومع ذلك، لا تتوافق هذه الاختبارات مع وصف تورينج. على سبيل المثال، نموذج باري (PARRY) الذي طوره كين كولبي خدع الأطباء النفسيين، وجعلهم يعتقدون أنهم يطلعون على مقابلات مع مصابين بجنون العظمة، حيث إنهم افترضوا بصورة طبيعية أنهم كانوا يتعاملون مع مرضى بشريين. وبالمثل، غالبًا ما يُنسب فن الكمبيوتر إلى الإنسان إذا لم يكن هناك إشارة على احتمالية استخدام الآلة.

أقرب شيء إلى اختبار تورينج الأصلي هو مسابقة لوبنر السنوية (التي تُعقد الآن في بلتشلي بارك). تنص القواعد على أن يكون هناك تفاعل لمدة ٢٥ دقيقة باستخدام

لكن هل هو ذكاء حقاً؟

٢٠ سؤالاً سبق وضعهم لاختبار الذاكرة والتفكير المنطقي والمعلومات العامة والجوانب الشخصية. تنظر هيئة التحكيم في مدى ملاءمة التعبير/القواعد وصحتها ووضوحها ومنطقيتها. وحتى الآن، لم يوجد برنامج قد خدع هيئة التحكيم في مسابقة لوبنر بمقدار ٣٠ بالمائة من الوقت. (في ٢٠١٤، قيل إن برنامجاً لطفل أوكراي يبلغ من العمر ١٣ عاماً خدع ٣٣ في المائة من المحققين، لكن الأخطاء تُغفر بسهولة عند غير الناطقين باللغة، لا سيما الأطفال).

مشكلات الوعي الكثيرة

لا يوجد شيء اسمه مشكلة الوعي. بل يوجد العديد من المشكلات الخاصة به. تُستخدم كلمة «الوعي» لتعيين العديد من الفروق مثل ما يأتي: اليقظة/النوم، التشاور/الغفلة، الانتباه/عدم الاكتراث، سهولة الوصول إلى الشيء/صعوبة الوصول إلى الشيء، التبليغ/عدم التبليغ، التفكير/عدم التفكير، وغيرها. لا يوجد شرح واحد يوضح كل تلك الفروق. الأضداد المدرجة عبارة عن أضداد وظيفية. وسيقول العديد من الفلاسفة إن تلك الأضداد قد تُفهم من حيث المبدأ في معالجة المعلومات ومصطلحات علم الأعصاب أو أحدهما.

لكن يبدو أن الوعي المدرك بالحواس مختلف، الأحاسيس (مثل الحزن أو الألم) أو الكيفيات المحسوسة (المصطلح الفني الذي يستخدمه الفلاسفة). ووجود الكيفيات المحسوسة في عالم مادي أساساً لغز ميتافيزيقي معروف أنه مستعصٍ. أطلق ديفيد تشالمرز على هذا اللغز «المسألة الصعبة». إنه يقول إنها مسألة لا مفر منها: «[يجب] أن نأخذ الوعي على محمل الجد ... وغير مقبول أن [نعيد] تعريف المسألة بأنها توضيح لطريقة أداء وظائف معرفية أو سلوكية معينة».

طُرحت العديد من الحلول ذات درجة التخمين العالية. تضمنت تلك الحلول إصدار تشالمرز عن النفسانية الشاملة، وهي نظرية تعترف أنها «شائنة أو حتى مجنونة»، وبموجبها يكون الوعي المدرك بالحواس خاصية في الكون لا يمكن اختزالها، وهي شبيهة بالكتلة أو الشحنة. لجأ العديد من أصحاب النظريات إلى فيزياء الكم، ويقول الخصوم إنهم يستخدمون لغزاً لحل لغز آخر. قال كولين ماكجين إن البشر غير قادرين دستورياً على فهم العلاقة بين العقل والكيفيات المحسوسة، مثلما لا تفهم الكلاب الحساب. كذلك يعتقد جيرى فودور، وهو من رواد فلسفة العلوم المعرفية، إذ يقول: «لا أحد لديه أدنى

فكرة عن كيف يمكن أن تعي المادة. ولا أحد يعلم حتى كيف سيكون الأمر إن كانت هناك أدنى فكرة عن كيف يمكن أن يعي أي شيء مادي.»
باختصار، يزعم قلة من الفلاسفة أنهم يفهمون الوعي المدرك بالحواس، ومن يزعمون ذلك لا يكاد يصدقهم أحد. فالموضوع شَرَك فلسفي.

وعي الآلة

المفكرون المتعاطفون مع الذكاء الاصطناعي يتعاملون مع الوعي بطريقتين. الطريقة الأولى بناء نماذج كمبيوتر خاصة بالوعي؛ ويطلق على تلك النماذج «وعي الآلة». والطريقة الثانية تحليل الوعي من منظور حاسوبي واسع من دون نمذجته (وهي سمة الفلاسفة المتأثرين بالذكاء الاصطناعي).

الذكاء الاصطناعي العام الذكي بحق سيكون له وعي وظيفي. على سبيل المثال، سيركز (ينتبه، يعي) الذكاء الاصطناعي على الأشياء المختلفة في أوقات مختلفة. ستتمكن الأنظمة التي ترتقي إلى مستوى الإنسان من التشاور والتفكر. بإمكانها أن تولد أفكاراً إبداعية، وأن تقيمها على نحو تشاوري. من دون تلك القدرات، فلا يمكن أن تولد أداءً ذكياً على ما يبدو.

يمكن أن يدخل الوعي المدرك بالحواس عندما يقيّم الإنسان الأفكار الإبداعية (انظر الفصل الثالث). في الحقيقة، يقول العديدون إنه داخل في كل فرق وظيفي. وعلى الرغم من ذلك، عادةً ما يتجاهل الباحثون في وعي الآلة الوعي المدرك بالحواس، حيث إنهم جميعاً يعتبرونه وعياً وظيفياً.

من مشاريع وعي الآلة المثيرة للاهتمام مشروع «تعلم وكيل التوزيع الذكي» الذي طورته مجموعة ستان فرانكلين في ممفيس. هذه التسمية تنطبق على شيئين. الأول هو نموذج مفاهيمي للوعي (الوظيفي)، وهو نظرية حاسوبية يعبر عنها بالألفاظ. والثاني عبارة عن تنفيذ جزئي وبسيط لذلك النموذج النظري.

وكلاهما يُستخدم لأغراض علمية (وهو الهدف الأساسي لفرانكلين). ولكن الشيء الآخر له تطبيقات عملية أيضاً. يمكن تخصيص تنفيذ مشروع تعلم وكيل التوزيع الذكي بحيث يتناسب مع نطاقات مسائل معينة، مثل الطب.

وعلى خلاف مشاريع التوجيه نحو النجاح وتحقيق الأهداف والتحكم المتكيف مع التفكير والعقلانية و«سي واي سي» (انظر الفصل الثاني)، فإن هذا المشروع حديث.

أُطلق الإصدار الأول عام ٢٠١١ (وقد صُمم للبحرية الأمريكية بهدف تنظيم الوظائف الجديدة للجنود الذين انتهت مدة خدمتهم). يتضمن الإصدار الحالي الانتباه وتأثيراته في التعلم في أنواع الذاكرة المتنوعة (العرضية والدلالية والإجرائية)، وفي الوقت الحالي ينفذ التحكم الحسي الحركي في الروبوتات. لكن لا تزال العديد من الميزات مفقودة، مثل اللغة. (الوصف التالي يتعلق بالنموذج المفاهيمي، بغض النظر عن الجوانب التي نُفذت بالفعل.)

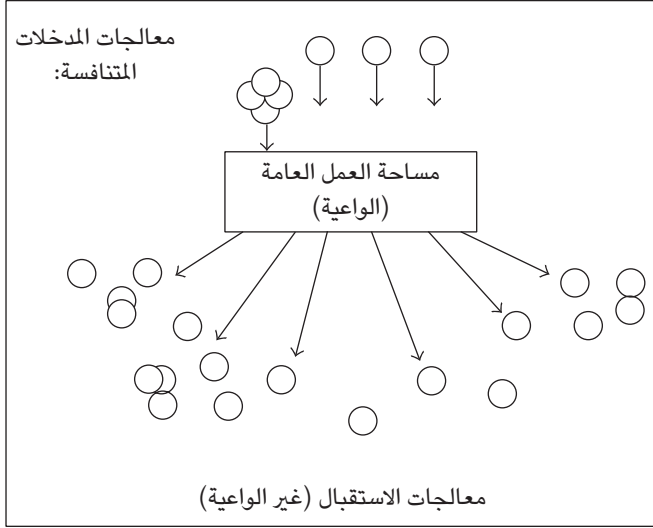
نظام «تعلم وكيل التوزيع الذكي» عبارة عن نظام مختلط؛ إذ يحتوي على تنشيط النشر والتمثيلات المنفردة (انظر الفصل الرابع)، وكذلك البرمجة الرمزية. إنه يعتمد على نظرية عصبية نفسية وضعها برنارد بارز، وهي نظرية مساحة العمل العامة ذات الصلة بالوعي.

ترى نظرية مساحة العمل العامة الدماغ على أنه نظام موزع (انظر الفصل الثاني)، يحتوي على مجموعة من الأنظمة الفرعية المتخصصة التي تعمل بالتوازي وتتنافس للوصول إلى الذاكرة قيد التشغيل (انظر شكل ٦-١). تظهر العناصر التي يحتوي عليها في ترتيب تسلسلي (تدفق الوعي)، ولكنها عبارة عن «محطة بث» لجميع المناطق القشرية. وإذا حفز عنصر بث — مستمد من عضو حسي أو نظام فرعي آخر — استجابة من منطقة قشرية معينة، فقد تتمتع تلك الاستجابة بالقوة الكافية للفوز بالمنافسة على الانتباه، وهو ما يتحكم في الوصول إلى الوعي تحكماً نشطاً. (عادةً ما تستحوذ الإدراكات/التمثيلات الجديدة على الانتباه، فالعناصر المتكررة تتلاشى من الوعي). غالباً ما تكون الأنظمة الفرعية معقدة. وبعض تلك الأنظمة متشابك هرمياً، والعديد منها له اتصالات ترابطية من أنواع مختلفة. تتشكل تجربة الوعي بمجموعة من السياقات اللاواعية (المنظمة في ذاكرات مختلفة)، وكلاهما يستحضر العناصر في مساحة العمل العامة ويعدلها. ومحتويات الانتباه بدورها تكيف السياقات الدائمة عن طريق التسبب في تعلم العديد من الأنواع.

عند بث هذه المحتويات، فإنها توجه اختيار الإجراء التالي. العديد من الإجراءات معرفية؛ بمعنى أنها تتمثل في بناء التمثيلات الداخلية أو تعديلها. كذلك تخزن المعايير الأخلاقية (في الذاكرة الدلالية) باعتبارها إجراءات لتقييم الإجراءات المحتملة. وأيضاً قد تتأثر القرارات بالتفاعلات المسبقة/المتوقعة من العوامل الاجتماعية الأخرى.

تأمل التخطيط على سبيل المثال (انظر الفصل الثاني). تمثل النوايا على أنها بنيات غير واعية إلى حد كبير، ولكنها عالية المستوى نسبياً، وقد تؤدي تلك البنيات إلى صور

الذكاء الاصطناعي



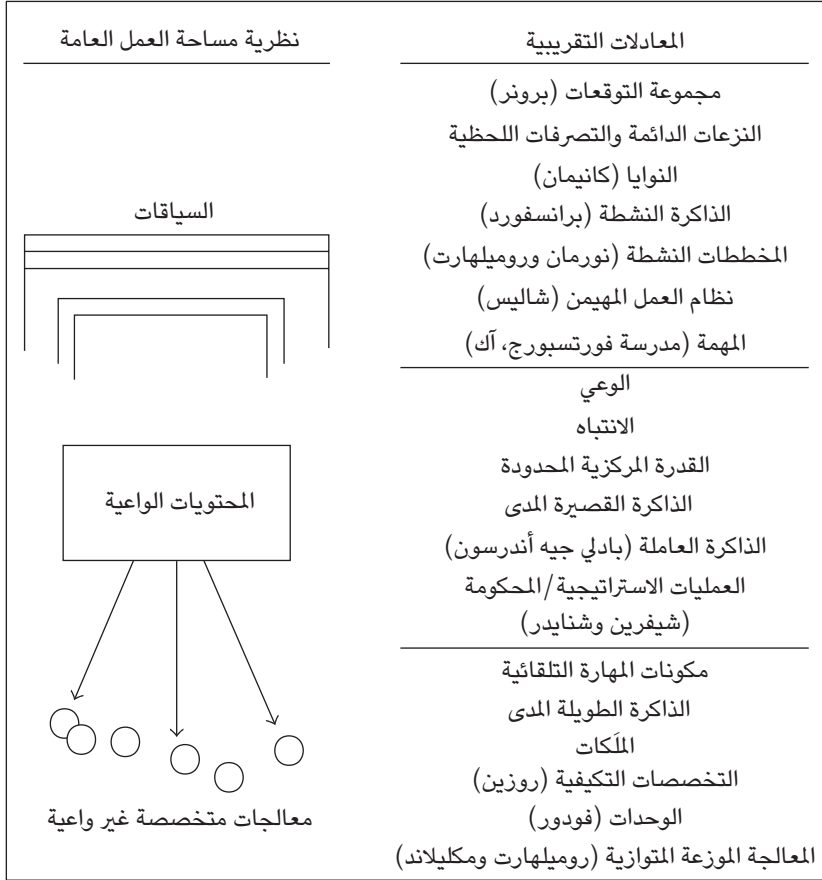
شكل ٦-١: مساحة العمل العامة في النظام الموزع. يحتوي الجهاز العصبي على معالجات متخصصة غير واعية (عوامل تحليل الإدراك، أنظمة مخرجات، أنظمة تخطيط، وغيرها). التفاعل بين هذه الأنظمة المتخصصة غير الواعية والتنسيق بينها والتحكم فيها يتطلب التبادل المركزي للمعلومات أو «مساحة العمل العامة». أدوات المدخلات المتخصصة يمكن أن تتعاون وتتنافس على الوصول إلى تلك المساحة. في الحالة الواردة هنا، تتعاون أربعة معالجات للمدخلات كي تترك رسالة عامة، وحينها تُبث الرسالة إلى النظام بأكمله.

أهداف واعية (وتختارها السمات البارزة في وقتها من الإدراك أو الذاكرة أو الخيال). وهذه توظف أهدافاً فرعية ذات صلة. إنها «توظف» ولا «تسترد»؛ حيث إن الأهداف الفرعية نفسها تقرر صلتها. ومثل كل الأنظمة الفرعية القشرية، فإنها تظل منتظرة حتى يحفزها عنصر بث؛ وفي هذا المقام يكون عبارة عن صورة هدف ملائمة. بإمكان نظام تعلّم وكيل التوزيع الذكي تحويل مخطط عمل موجه بالأهداف إلى إجراءات حركية منخفضة المستوى يمكن تنفيذها، وتستجيب إلى السمات التفصيلية في بيئة غير متوقعة ومتغيرة.

نظرية بارز (وإصدار فرانكلين منها) لم يتخيلها أحد في أوساط علماء الكمبيوتر. على النقيض من ذلك، وُضعت النظرية كي تأخذ في الاعتبار مجموعة كبيرة من الظواهر

لكن هل هو ذكاء حقاً؟

النفسية المعروفة، ومجموعة كبيرة من الأدلة التجريبية (انظر شكل ٦-٢). لكن يزعم المؤلفان أن هذه النظرية تحل بعض الألغاز النفسية التي لم يسبق حلها.



شكل ٦-٢: أوجه التشابه بين مساحة العمل العامة وغيرها من المفاهيم المنتشرة. كل فكرة من هذه الأفكار المألوفة تعرفها (نظرية مساحة العمل العامة) من حيث الوظائف غير الواعية والوظائف الواعية.

على سبيل المثال، يقولان إن نظرية مساحة العمل العامة/نظام تعلّم وكيل التوزيع الذكي تحل مشكلة «الربط» التي يدور حولها جدال منذ فترة طويلة. هذا يطرح سؤالاً، وهو كيف لمدخلات متنوعة من حواس مختلفة في مناطق مختلفة من الدماغ — مثل الشعور والشكل الخارجي والشم والصوت لدى قطة — تُعزى إلى شيء واحد فقط. يقول فرانكلين وبارز إن هذه النظرية تشرح أيضاً كيف أن الدماغ البشري يتجنب مشكلة الإطار (انظر الفصل الثاني). عند توليد أمثلة إبداعية على سبيل المثال، فإنه لا توجد أداة تنفيذ مركزية، بل يُبحث في بنية البيانات بكاملها عن العناصر ذات الصلة. بل إذا أدرك النظام الفرعي أن عنصر بـ يتلاءم/يقترّب من الشيء الذي يبحث عنه (دائماً)، فإنه يتنافس للدخول إلى مساحة العمل العامة.

يذكّرنا منهج الذكاء الاصطناعي ذاك بـ «البرامج الخفية» في نظام بانديمونوم، والبنيات المعمارية لـ «الألواح» المستخدمة في تنفيذ أنظمة الإنتاج (انظر الفصلين الأول والثاني). هذا ليس مدهشاً؛ حيث إن تلك الأفكار ألهمت نظرية بارز النفسية العصبية التي أدت في النهاية إلى إنشاء نظام تعلّم وكيل التوزيع الذكي. عادت عجلة النظريات إلى نقطة البداية.

الذكاء الاصطناعي والوعي المدرك بالحواس

يتجاهل ممارسو وعي الآلة المشكلة الصعبة. لكن ثلاثة فلاسفة استلهموا الحل من الذكاء الاصطناعي وحلّوها مباشرةً، وهم: بول تشيرشلاند ودانييل دينيت وآرون سلومان. والقول بأن إجاباتهم مثيرة للجدل قولٌ بخس. لكن بما أن الأمر له علاقة بالوعي المدرك بالحواس، فهذا يتسق مع التوقعات.

تنكر «المادية الإقصائية» لتشيرشلاند وجود الأفكار والتجارب غير المادية. بل إنها تعرّفها بأنها حالات عقلية.

إنه يطرح نظريةً جزء منها حاسوبي (ترابطي) وجزء آخر عصبي، وتلك النظرية تحدد «مساحة تذوق» رباعية الأبعاد تتسم بالمنهجية في وضع فروق موضوعية (الكيفيات المحسوسة) للتذوق في بنيات عصبية محددة. الأبعاد الأربعة تعبر عن المستقبليات الأربعة الخاصة بالتذوق في اللسان.

بالنسبة إلى تشيرشلاند، تلك المسألة ليست علاقة بين العقل والدماغ؛ الحصول على تجربة تذوق يعني ببساطة أن يزور الدماغ نقطة معينة في تلك المساحة الحسية المحددة

لكن هل هو ذكاء حقاً؟

تحديداً مجرداً. المعنى الضمني هو أن كل الوعي المدرك بالحواس عبارة عن وجود الدماغ في موقع محدد في مساحة تشعبية يمكن اكتشافها بالتجربة. وإذا كان الأمر كذلك، فعندئذٍ لن يكون هناك جهاز كمبيوتر (ربما باستثناء محاكاة الدماغ بالكامل) يتمتع بالوعي المدرك بالحواس.

ينكر دينيت أيضاً وجود تجارب مميزة وجودياً، فضلاً عن الأحداث الجسدية. (لذا، الرد المشهور على كتابه الاستفزازي الصادر عام ١٩٩١ هو: «لم يوضَّح الوعي، ولكن انتُحلت له الأعذار»).

يرى أن التجربة هي التمييز. لكن في تمييز الأشياء الموجودة في العالم المادي، فإن المرء لا يأتي بشيء آخر إلى الوجود في عالم آخر غير مادي. وقد عبّر عن ذلك في محادثة من خياله:

[أوتو]: يبدو لي أنك تنكر وجود معظم الظواهر الحقيقية التي لا شك فيها، وتلك المظاهر لم ينكرها حتى ديكارت في كتابه «التأملات».

[دينيت]: بمعنى أو بآخر، أنت محقٌّ، هذا ما أنكر وجوده. [لنفكر] في ظاهرة انتشار لون النيون. يبدو أن هناك حلقة متوهجة بلون مائل إلى الوردى على الغلاف الورقي للكتاب. [إنه يصف سراباً بصرياً تسببه الخطوط الحمراء والسوداء على ورقة بيضاء لامعة].

[أوتو]: لا شك في وجود تلك الظاهرة.

[دينيت]: لكن لا يوجد أي حلقة مائلة إلى اللون الوردى. ليس على وجه الحقيقة.

[أوتو]: صحيح. ولكن يبدو أن هناك حلقة وردية بلا شك!

[دينيت]: صحيح.

[أوتو]: إذن، أين هي؟

[دينيت]: أين ماذا؟

[أوتو]: الحلقة المتوهجة المائلة إلى اللون الوردى.

[دينيت]: لا توجد أي حلقات، اعتقدت أنك اعترفت بذلك للتو.

[أوتو]: نعم، لا توجد أي حلقات مائلة إلى اللون الوردى على الورقة، ولكن لا شك

أنها موجودة.

[دينيت]: صحيح. يبدو أن هناك حلقة متوهجة مائلة إلى اللون الوردى.

[أوتو]: إذن، لننتحدث عن تلك الحلقة.

[دينيت]: أي واحدة؟

[أوتو]: التي يبدو أنها موجودة.

[دينيت]: لا يوجد شيء في شكل حلقة وردية وتبدو وكأنها موجودة.

[أوتو]: اسمع، أنا لا أقول يبدو أن هناك حلقة متوهجة مائلة إلى اللون الوردى، بل

يبدو أن هناك حلقة متوهجة مائلة إلى اللون الوردى بالفعل!

[دينيت]: أنا أتسرع في موافقتك. أنت تعني ما تقوله بالفعل حين تقول إنه يبدو

أن هناك حلقة متوهجة مائلة إلى اللون الوردى.

[أوتو]: اسمع. أنا لا أقصد. أنا لا أعتقد فقط أن هناك حلقة متوهجة مائلة إلى اللون

الوردى، بل هناك حلقة متوهجة مائلة إلى اللون الوردى بالفعل!

[دينيت]: لقد فعلتها الآن. لقد وقعت في فخ كما وقع فيه كثيرون غيرك. وكأنك

تعتقد أنه لا يوجد فرق بين الاعتقاد (إصدار حكم، اتخاذ القرار، التثبت من رأي) بأن

شيئاً كأنه وردى في نظرك و شيئاً تراه «بالفعل» بلون وردى. لكن لا يوجد فرق. لا توجد

ظاهرة تُرى وكأنها حقيقة، بالإضافة إلى ظاهرة الحكم بطريقة أو بأخرى بأن ذلك الشيء هو الواقع.

بعبارة أخرى، لا يمكن الوفاء بمتطلبات شرح الكيفيات المحسوسة. فهذه الأشياء

ليست موجودة.

لا يتفق آرون سلومان مع هذا الرأي. إنه يُقرُّ بالوجود الفعلي للكيفيات المحسوسة.

لكنه يُقرُّ بطريقة غير عادية؛ إنه يحلُّها باعتبارها أنماطاً لأجهزة افتراضية متعددة

الأبعاد، ونحن نطلق عليها اسم العقل (انظر القسم التالي).

يقول إن الكيفيات المحسوسة عبارة عن حالات حاسوبية داخلية. وقد يكون لها

تأثيرات سببية في السلوك (مثل تعبيرات الوجه اللاإرادية) وفي الأنماط الأخرى لمعالجة

المعلومات في العقل. إنها لا توجد إلا في الأجهزة الافتراضية ذات التعقيد الهيكلي الكبير

(إنه يوضح أنواع الموارد الحاسوبية المنعكسة المطلوبة). لا يمكن الوصول إليها إلا ببعض

الأجزاء الأخرى في الجهاز الظاهري المعنى، وليس بالضرورة أن ينطوي على أي تعبير

سلوكي. (ومن هنا تأتي الخصوصية). إضافة إلى ذلك، لا يمكن وصفها بعبارات لفظية

على الدوام باستخدام مستويات العقل العليا الذاتية الرصد. (ومن هنا يأتي عدم إمكانية

وصفها).

هذا لا يعني أن سلومان يحدد الكيفيات المحسوسة بعمليات الدماغ (مثلاً يفعل

تشيرلاندر). فالحالات الحاسوبية عبارة عن أنماط للأجهزة الافتراضية؛ ومن ثم لا يمكن

لكن هل هو ذكاء حقاً؟

تحديدها بلغة الأوصاف المادية. ولكن لا يمكن أن توجد ويكون لها تأثيرات سببية إلا عند تنفيذها ببعض الآليات المادية الكامنة.

ماذا عن اختبار تورينج؟ تقول تحليلات كلٍّ من دينيت وسلومان (ويقول دينيت صراحةً) بشكلٍ ضمني إن مفاهيم الزومبي غير ممكنة. والسبب في رأيهم أن مفهوم «الزومبي» غير مترابط. بتوافر السلوك المناسب والجهاز الظاهري أو أحدهما، يرى سلومان الوعي بأنه مضمون حتى وإن كان يتضمن الكيفيات المحسوسة. ومن ثم يُحفظ اختبار تورينج من الاعتراض بأنه يمكن أن «يجتازه» زومبي.

وماذا عن الذكاء الاصطناعي العام الافتراضي؟ إذا كان دينيت على صواب، فسيتضمن رأيه كل الوعي الذي نعيه، ولكن من دون الكيفيات المحسوسة. وإذا كان سلومان على صواب، فسيتضمن رأيه الوعي المدرك بالحواس بالمعنى نفسه الذي لدينا.

الأجهزة الافتراضية ومعضلة العقل والجسم

طرح هيلاري بوتنام نظرية «الوظيفية» في ستينيات القرن العشرين، استخدمت النظرية فكرة آلات تورينج، وكذلك التمييز بين البرمجيات والأجهزة (المستحدثة حينذاك) كي يقول — فعلياً — إن العقل هو وظيفة الدماغ.

الانقسام الميتافيزيقي (الديكارتي) بين مادتين مختلفتين تماماً، أفسح الطريق إلى الانقسام المفاهيمي بين مستويات الوصف. المقارنة بين البرنامج والكمبيوتر أفادت بأن «العقل» مختلفٌ تمام الاختلاف عن «الجسم». ولكن تلك الفكرة تتماشى تماماً مع المادية. (ثار جدلٌ كبير بشأن ما إذا كانت الفكرة تنطوي على الكيفيات المحسوسة، ولا يزال الجدل قائماً).

على الرغم من أن العديد من برامج الذكاء الاصطناعي المثيرة للاهتمام طُورت بحلول عام ١٩٦٠ (انظر الفصل الأول)، فنادرًا ما فُكّر الفلاسفة من أصحاب المذهب الوظيفي في أمثلة محدّدة. فقد ركّزوا على المفاهيم الأساسية، مثل حوسبة تورينج. وبمجرد ظهور المعالجة الموزعة المتوازية (انظر الفصل الرابع) في منتصف ثمانينيات القرن العشرين، فُكّر العديد من الفلاسفة بالفعل في طريقة عمل أنظمة الذكاء الاصطناعي. وحتى ذلك الحين، لم يتساءل سوى قلةٍ قليلةٍ ما الوظائف الحاسوبية التي يمكن أن تجعل التفكير أو الإبداع (على سبيل المثال) ممكنًا.

أفضل طريقة لفهم تلك المسائل هي اقتراض مفهوم الأجهزة الافتراضية الذي صاغه علماء الكمبيوتر. بدلاً من القول بأن العقل هو وظيفة الدماغ، ينبغي القول (تبعاً لسلومان)

بأن العقل عبارة عن جهاز ظاهري، أو بالأحرى مجموعة متكاملة تتكون من عدة أجهزة افتراضية مغروسة في الدماغ. (ولكن مطابقة العقل بالجهاز الظاهري له تأثير مناقض جداً للبدئية؛ انظر قسم «هل البروتين العصبي ضروري؟»).

حسب ما ورد في الفصل الأول، الأجهزة الافتراضية حقيقة ولها تأثيرات سببية؛ بمعنى أنه لا يوجد تفاعلات ميتافيزيقية غامضة بين العقل والجسد. ومن ثم تكمن الأهمية الفلسفية لمشروع تعلّم وكيل التوزيع الذكي — على سبيل المثال — في تحديد مجموعة منظمة من الأجهزة الافتراضية التي تبين مدى احتمالية الأنماط المتنوعة للوعي (الوظيفي). يعدّل نهج الأجهزة الافتراضية نمطاً أساسياً لمبدأ الوظيفية، وهو فرضية نظام الرموز الفيزيائية. في سبعينيات القرن العشرين، عرّف ألين نيويل وهيربرت سيمون فرضية نظام الرموز الفيزيائية بأنها «مجموعة كيانات تُسمى الرموز، وهي عبارة عن أنماط فيزيائية يمكن أن تكون في صورة مكونات لنوع آخر من الكيانات يُسمى التعبير (أو هيكل الرمز). ... [داخل] بنية الرموز ... ترتبط مثيلات الرموز (أو الرموز المميزة) بطريقة فيزيائية (كأن يكون رمز مميز بجوار رمز مميز آخر)». يقولون تهدف عمليات المعالجة إلى إنشاء بنيات الرموز وتعديلها؛ بمعنى عمليات المعالجة التي يحددها الذكاء الاصطناعي الرمزي. وقالوا أيضاً «إن فرضية نظام الرموز الفيزيائية تنطوي على الوسائل الضرورية والكافية للإجراء الذكي العام». بعبارة أخرى، العلاقة بين العقل والدماغ عبارة عن نظام رموز فيزيائية.

ومن منظور مطابقة العقل بالجهاز الظاهري، كان ينبغي أن يُطلق على تلك الفرضية «نظام الرموز المطبق فيزيائياً» (لنتجنب التعبير عن ذلك في صورة اختصار)، حيث إن الرموز هي محتويات الأجهزة الافتراضية وليس الأجهزة المادية. هذا يعني أن الأنسجة العصبية ليست ضرورية للذكاء ما لم تكن المادة الأساسية الوحيدة القادرة على تنفيذ الأجهزة الافتراضية المعنية.

تقترح افتراضية نظام الرموز الفيزيائية (وكثير من أنظمة الذكاء الاصطناعي الأولى) أن التمثيل — أو الرموز الفيزيائية — سمة يمكن عزلها بوضوح أو تحديد موقعها بدقة في الجهاز/الدماغ. ستتقدم الترابطية معنًى مختلفاً تماماً عن التمثيلات (انظر الفصل الرابع). إنها ترى التمثيلات من حيث شبكات الخلايا بالكامل، وليس الخلايا العصبية التي يمكن تحديد مواقعها بوضوح. كذلك رأت المفاهيم من حيث القيود المتضاربة جزئياً، وليس التعريفات المنطقية الصارمة. وتلك الفكرة استهوت إلى حد كبير الفلاسفة المطلعين على وصف لودفيج فيتجنشتاين عن «التشابهات العائلية».

بعد ذلك، أنكر العاملون في الروبوتات الكائنة أن الدماغ يحتوي على تمثيلات على الإطلاق (انظر الفصل الخامس). هذا الموقف قبله بعض الفلاسفة، ولكن ديفيد كيرش على سبيل المثال احتجَّ بأن التمثيلات التركيبية (والحوسبة الرمزية) ضرورية لكل السلوكيات التي تحتوي على مفاهيم، بما في ذلك المنطق واللغة والتشاورات.

القصد والفهم

طبقاً لنيويل وسيمون، أي نظام رموز فيزيائية ينفذ العمليات الحاسوبية الصحيحة ذكي بالفعل. إنه يتضمن «الوسائل الضرورية والكافية للتصرف بذكاء». يطلق الفيلسوف جون سيرل على هذا المطلب «الذكاء الاصطناعي القوي». («الذكاء الاصطناعي الضعيف» يرى أن نماذج الذكاء الاصطناعي تساعد فقط علماء النفس على صياغة نظريات متسقة). يقول إن الذكاء الاصطناعي القوي كان مخطئاً. ربما تقبّع الحوسبة الرمزية في رءوسنا (على الرغم من أنه يشك في ذلك)، ولكن لا يمكنها وحدها أن توفر الذكاء. بمزيد من الدقة، إنها لا توفر «القصدية»، وهي اللفظة التي اصطلح عليها الفلاسفة، وتعني القصد والفهم.

اعتمد سيرل على تجربة فكرية لا تزال محل جدل حتى اليوم، وهي: «يوجد سيرل في غرفة من دون نوافذ، ويوجد بها فتحة تُمرَّر منها قصاصات ورق مكتوب عليها رموز عبثية. يوجد صندوق من القصاصات المرسوم عليها رموز عبثية مماثلة، ومعها كتاب قواعد يقول إذا مرَّر رمز عبثي فعلى سيرل أن يُخرج رمزاً عبثياً، أو ربما يطلع على سلسلة طويلة من أزواج الرموز العبثية قبل أن يُخرج قصاصة من الفتحة. ومن دون أن يعلم سيرل، كانت الرسوم العبثية كتابة صينية، وكتاب القواعد عبارة عن برنامج معالجة اللغة الطبيعية الصينية، والصينيون خارج الغرفة يستخدمون سيرل للرد على أسئلتهم. على الرغم من ذلك، دخل سيرل الغرفة من دون أن يكون قادراً على فهم اللغة الصينية، وسيبقى غير قادر على فهمها عندما يغادر الغرفة». الاستنتاج: «الحوسبة الصورية وحدها (وهي ما يفعله سيرل داخل الغرفة) لن تتولد عنها القصدية. ومن ثم، فالذكاء الاصطناعي القوي مخطئ، ويستحيل تحصيل الفهم الحقيقي في برامج الذكاء الاصطناعي.» (كان الذكاء الاصطناعي الرمزي هو الهدف الأساسي من تلك الحجة التي تُسمى «الغرفة الصينية»، ولكنها عُمت بعد ذلك وانطبقت على الترابطية وعلم الروبوتات).

زعم سيرل هنا مُفاده أن «المعاني» المنسوبة إلى برامج الذكاء الاصطناعي تنبع بالكامل من المستخدم البشري/المبرمج. تلك المعاني مطلقة فيما يتعلق بالبرنامج نفسه،

حيث إنه فارغ من حيث الدلالات. وبما أن البرنامج عبارة عن «تراكيب لغوية من دون دلالات»، فإنه يمكن تفسير البرنامج نفسه على أنه جدول حسابات ضرائب أو تصاميم حركات.

وهذا صحيح في بعض الأحيان. ولكن تذكر أن مطالبة فرانكلين بأن نماذج تعلم وكيل التوزيع الذكي رُسِّخت — بل جسَّدت — المعرفة باستخدام وسائل الاقتران المنظم بين المعاني ووحدات التشغيل والبيئة. تذكر أيضًا وحدة التحكم التي تطورت وأصبحت وحدة اكتشاف الاتجاهات لدى الروبوت (انظر الفصل الخامس). وإطلاق اسم «وحدة اكتشاف الاتجاهات» ليس اعتباطيًا. فوجودها يعتمد على تطورها باعتبارها وحدة اكتشاف اتجاهات مفيدة في تحقيق هدف الروبوت.

المثال الأخير وثيق الصلة، لا سيما أن بعض الفلاسفة يرون التطور على أنه مصدر القصدية. تقول روث ميلكان على سبيل المثال إن التفكير واللغة من الظواهر البيولوجية، وتعتمد معانيهما على التاريخ التطوري. وإذا صح ذلك، فلا يمكن أن يكون للذكاء الاصطناعي العام غير التطوري فهم حقيقي.

الفلاسفة الآخرون ذوو التفكير العلمي (مثل نيويل وسيمون أنفسهم) يعرفون القصدية من منظور السببية. لكنهم يواجهون صعوبة في تفسير العبارات غير الصحيحة: إذا ادَّعى شخص ما رؤية بقرة، ولكن لا توجد بقرة هناك تدفعه للتلفظ بتلك الكلمات، فما الذي تعنيه كلمة بقرة؟

بإيجاز، لا توجد نظرية عن القصدية تتفق مع آراء كل الفلاسفة. ولأن الذكاء الحقيقي يتضمن الفهم، فهذا سبب آخر لعدم معرفة أي شخص ما إذا كان الذكاء الاصطناعي العام الافتراضي لدينا سيكون ذكيًا حقًا أم لا.

هل البروتين العصبي ضروري؟

من الأسباب التي رفض سيرل الذكاء الاصطناعي القوي من أجلها، أن أجهزة الكمبيوتر ليست مصنوعة من البروتين العصبي. قال إن القصدية ناتجة عن البروتين العصبي مثلما ينتج البناء الضوئي عن الكلوروفيل. قد لا يكون البروتين العصبي المادة الوحيدة في الكون التي بإمكانها دعم القصدية والوعي. لكن قال إنه من الواضح أن المعدن والسيليكون لا يدعمانها.

لكن هل هو ذكاء حقاً؟

تلك الخطوة بعيدة للغاية. من المسلّم به أنه من غير البديهي للغاية أن نقول إن أجهزة الكمبيوتر المصنوعة من القصدير يمكن أن تعاني حقاً من الحزن أو الألم أو تفهم اللغة حقاً. ولكن الكيفيات المحسوسة الناتجة عن البروتين العصبي لا تقل في منافاتها للبديهة وإثارة الإشكاليات الفلسفية. (ومع ذلك، فإن الشيء المنافي للبديهة قد يكون صحيحاً).

إذا قبل أحد تحليل الكيفيات المحسوسة الذي أجراه جهاز سلومان الظاهري، فستتلاشى هذه الصعوبة الخاصة. لكن التصوير العام للعقل بالجهاز الظاهري يستحضر صعوبة أخرى مماثلة. إذا نُفِذَ جهاز ظاهري مؤهل للعقل في جهاز يعمل بالذكاء الاصطناعي، فإن ذلك العقل سيوجد في الجهاز، أو ربما في عدة أجهزة. لذا، فإن مطابقة العقل بالجهاز الظاهري يعني من حيث المبدأ إمكانية خلود الأشخاص (مضاعفة الاستنساخ) في أجهزة الكمبيوتر. ومن وجهة نظر الناس (انظر الفصل السابع)، هذا لا يقل في منافاته للبديهة عن أجهزة الكمبيوتر التي تدعم الكيفيات المحسوسة.

إذا كان البروتين العصبي هو المادة الوحيدة القادرة على دعم الأجهزة الافتراضية على مستوى الإنسان، فيمكننا نبذ مقترح «الخلود المستنسخ». لكن هل الأمر كذلك؟ لا نعلم. ربما يحتوي البروتين العصبي على خصائص خاصة، أو ربما ذات درجة تجريد عالية، وهذه الخصائص تمكّنه من تنفيذ هذه المجموعة الكبيرة من العمليات الحاسوبية التي ينفذها العقل. على سبيل المثال، يجب أن يمتلك القدرة على استخلاص الجزئيات المستقرة (بسرعة كبيرة نوعاً ما)، وكذلك (القابلة للتخزين) والمرنة أيضاً. يجب أن يكون قادراً على تكوين البنيات والاتصالات بين تلك البنيات التي تحتوي على خصائص كهروكيميائية تمكّنها من تمرير المعلومات فيما بينها. وربما توجد مواد أخرى على كواكب أخرى تؤدي تلك المهمات أيضاً.

ليس الدماغ فحسب، بل الجسم أيضاً

يقول بعض فلاسفة العقل إن الدماغ يحظى بقدر كبير من الاهتمام. يقولون إنه من الأفضل أن ينصبّ التركيز على الجسد بكامله.

يستند موقفهم على الظواهر القارية التي تؤكد على «شكل الحياة» الإنسانية. وهذا يشمل كلاً من الوعي الهادف (بما في ذلك «مصالح» الإنسان التي ينبني عليها إحساسنا بالصلة بالموضوع) والتجسيد.

التجسيد يعني أن يُوجد الإنسان في بيئة ديناميكية ويتفاعل معها تفاعلاً نشطاً. تعني البيئة — والتفاعل — كلاً من العوامل الفيزيائية والاجتماعية الثقافية. الخصائص النفسية ليست التفكير المنطقي والأفكار، ولكن التكيف والتواصل.

لا يولي فلاسفة التجسيد وقتاً طويلاً للذكاء الاصطناعي الرمزي؛ فهم يرونه يهتم بجانب العقل اهتماماً مُفرطاً. وهم لا يفضلون غير النُهج القائمة على السبرانية (انظر الفصلين الأول والخامس). وبما أن الذكاء الحقيقي يعتمد على الجسم وفقاً لوجهة النظر هذه، فلا يمكن أن يكون الذكاء الاصطناعي على الشاشة ذكياً حقاً. حتى إذا كان النظام على الشاشة عبارة عن عامل مستقل هيكلياً ومقترن ببيئة مادية، فلن يُعتبر (كما هي الحال مع فرانكلين) مجسداً.

ماذا عن الروبوتات؟ في النهاية، الروبوتات عبارة عن كائنات مادية توجد في عالم حقيقي وتتكيف معه. وفي الحقيقة، يُثني هؤلاء الفلاسفة على الروبوتات الكائنة في بعض الأحيان. لكن هل الروبوتات لها أجسام؟ أو هل لها مصالح؟ أو هل تشكل حياة؟ هل هم أحياء على الإطلاق؟

سيجيب علماء الظواهر «بالتأكيد لا!» ربما يستشهدون بملاحظة فيتجنشتاين الشهيرة: «إذا كان بإمكان الأسد أن يتكلم، فلن نفهمه.» نمط حياة الأسد مختلف عن نمط حياة البشر لدرجة أن التواصل يكاد يكون مستحيلاً. لا شك أن هناك تداخلاً بين الجوانب النفسية لدى الأسد ومثيلاتها لدى الإنسان (مثل الجوع، والخوف، والتعب، وغير ذلك) لدرجة أنه قد يجدي قدر ضئيل من الفهم والتعاطف. ولكن حتى هذا لن يتوفر عند «التواصل» مع الروبوت. (لذلك، الأبحاث على أجهزة الكمبيوتر الرفيعة تثير المخاوف؛ انظر الفصلين الثالث والسابع).

مجتمع أخلاقي

هل سنقبل — أو ينبغي أن نقبل — بوجود أجهزة الذكاء الاصطناعي العام بمستوى ذكاء الإنسان باعتبارها أعضاء في مجتمعنا الأخلاقي؟ وإن قبلناها، فسيكون لهذا تبعات عملية خطيرة. فتلك الأنظمة ستؤثر في الإنسان؛ حيث إن الكمبيوتر يتفاعل بثلاث طرق. الطريقة الأولى: يستقبل الذكاء الاصطناعي العام اهتمامنا الأخلاقي، مثلما تفعل الحيوانات. ونحن سنحترم مصالحه إلى حد معين. وإذا طُلب من أحد أن يقطع وقت راحته أو حل لغز الكلمات المتقاطعة كي يساعده في تحقيق هدف «له أولوية عليا»،

لكن هل هو ذكاء حقاً؟

فسيفعل الشخص ذلك. (ألم يسبق ونهضت من على كرسيك كي تتمشى بالكلب أو تدع الدعسوقة تخرج إلى الحديقة؟) كلما ارتفع مستوى حُكمنا بأن مصالحهم تهمنا، زادت درجة التزامنا باحترامها. لكن قد يعتمد الحكم إلى حد كبير على ما إذا كنا نسند الوعي المدرك بالحواس (بما في ذلك العواطف المحسوسة) إلى الذكاء الاصطناعي العام أم لا.

الطريقة الثانية: سنعتبر أفعال تلك الأنظمة قابلة للتقييم الأخلاقي. الطائرات بدون طيار القاتلة لا تتحمل مسؤولية أخلاقية (على خلاف مستخدميها/مصمميها؛ انظر الفصل السابع). لكن هل يتحمل نظام الذكاء الاصطناعي العام الذكي بحق المسؤولية؟ من المفترض أن تتأثر قرارات النظام بتفاعل الإنسان تجاهها؛ بالمدح أو الذم من جانب الإنسان. وإذا لم يكن كذلك، فلن يكون هناك مجتمع. بإمكان النظام أن يتعلم «الأخلاقيات» بقدر ما يستطيع الرضيع (أو الكلب) أن يكون حسن السلوك، أو الطفل الكبير أن يكون مراعيًا لرغبات الآخرين. (تتطلب مراعاة الآخرين تطوير ما يُسميه علماء النفس «نظرية العقل» التي تفسر سلوك الإنسان من حيث القدرة على الفعل والنية والاعتقاد). حتى العقاب يمكن تبريره بموجب مسوغات جوهرية.

الطريقة الثالثة: يمكن أن نجعل النظام هدفًا للحجة والإقناع بشأن القرارات الأخلاقية. يمكن كذلك أن يقدم مشورات أخلاقية للناس. كي نخرط بجدية في هذه المحادثات، فيجب أن نكون على ثقة بأن النظام يقبل الاعتبارات الأخلاقية على وجه التحديد (بالإضافة إلى التمتع بالذكاء على مستوى ذكاء الإنسان). ولكن ما الذي يعنيه ذلك؟ لا يوافق أصحاب فلسفة الأخلاق بشدة على محتوى الأخلاقيات فحسب، بل لا يوافقون على أساسها الفلسفي.

كلما نظر المرء في الآثار المترتبة على «المجتمع الأخلاقي»، بدا أن فكرة الاعتراف بالذكاء الاصطناعي العام تزيد إشكالياتها. في الواقع، لدى معظم الناس حدس قوي بأن الاقتراح نفسه سخيف.

الأخلاقيات والحرية والنفس

ينشأ هذا الحدس إلى حد كبير لأن مفهوم المسؤولية الأخلاقية يرتبط ارتباطاً وثيقاً بمفاهيم أخرى — القدرة الواعية على الفعل والحرية والنفس — ومن ثم يساهم في فكرتنا عن الإنسانية.

المشاورات الواعية تزيد مقدار المسؤولية الأخلاقية في اختياراتنا (على الرغم من إمكانية انتقاد الأفعال غير المدروسة أيضاً). يُعزى المدح أو الذم الأخلاقي إلى الفاعل أو

النفس التي ارتكبت الفعل. كذلك الأفعال التي ارتكبت في ظل قيود صارمة أقل عرضة للذم أكثر من الأفعال التي ارتكبت بحرية.

تثير هذه المفاهيم كثيرًا من الجدل حتى عندما تنطبق على الناس. لا يبدو من الملائم تطبيقها على الآلات لا سيما أنها ناتجة عن آثار التفاعل بين الإنسان والكمبيوتر التي استشهدنا بها في القسم السابق. وعلى الرغم من ذلك، فإن اتباع منهج «مطابقة العقل بالجهاز الظاهري» وتطبيقه على العقول البشرية هو أمرٌ يمكن أن يساعدنا على فهم هذه الظاهرة في حالتنا نحن.

يحلل الفلاسفة المتأثرون بالذكاء الاصطناعي الحرية من منظور أنواع معينة من التعقيد التحفيزي المعرفي. إنهم يشيرون إلى أنه لا يخفى أن الإنسان يتمتع بـ «حريات» لا تتمتع بها صراصير الليل على سبيل المثال. تجد أنثى صرصور الليل وليفها عن طريق استجابات منعكسة مرتبطة بعضها ببعض (انظر الفصل الخامس). لكن المرأة التي تبحث عن زوج تتوافر لها العديد من الاستراتيجيات. كذلك يكون لها أهداف أخرى غير الزواج، ولا يمكن إشباعها كلها في وقتٍ واحد. وعلى الرغم من ذلك، فإنها تستطيع فعل أشياء لا تستطيعها الصراصير بفضل الموارد الحاسوبية (أي الذكاء).

تأصلت هذه الموارد بالوعي الوظيفي، وتتضمن التعلم الإدراكي والتخطيط المسبق والواجب الافتراضي وترتيب التفضيلات والتفكير في الواقع المضاد وجدولة الأفعال الموجهة بالعاطفة. وفي الواقع، يستخدم دينيت في كتابه «المجال الفسيح» تلك المفاهيم لشرح حرية الإنسان، ويضرب مجموعة من الأمثال عليها. ومن ثم يساعدنا الذكاء الاصطناعي في فهم إلى أي مدى نتمتع بالحرية في الاختيار.

الجبرية/التخيير مسألة مضللة إلى حدٍ كبير. يوجد قدر من التخيير في أفعال البشر، ولكن هذا لا يحدث عند الوصول إلى القرار؛ لأن هذا لن يخيره في المسؤولية الأخلاقية. ومع ذلك، قد يؤثر التخيير في الاعتبارات التي تنشأ في أثناء المشاورات. قد يفكر العامل في القرار «س» أو يذكر بالقرار «ص»، حيث إن «س» و«ص» يتضمنان كلاً من الحقائق والقيم الأخلاقية. على سبيل المثال، اختيار أحدٍ ما هدية عيد الميلاد قد يتأثر بملاحظة عرضية لشيء ما يذكره بأن مستلم الهدية يحب اللون الأرجواني أو يدعم حقوق الحيوان. كل الموارد الحاسوبية المدرجة للتو ستتوفر للذكاء الاصطناعي العام على مستوى الإنسان. ولذا، إذا لم يتوجب أن يتضمن الاختيار الحر الوعي المدرك بالحواس (وإذا لم يرفض أحد التحليلات الحاسوبية الخاصة بذلك الاختيار)، فيبدو أن الذكاء الاصطناعي

لكن هل هو ذكاء حقاً؟

العام الذي نتخيله سيحظى بالحرية. حتى إن كنا نفهم أن للذكاء الاصطناعي العام دوافع متنوعة تهّمه، فإنه يمكن التمييز بين الاختيارات التي اتخذها بـ «حريته» أو التي اتخذها «تحت قيود». لكن ذلك «الشرط» واسع للغاية.

بالنسبة إلى النفس، يؤكد الباحثون في الذكاء الاصطناعي على دور الحوسبة التكرارية، وفي ذلك يمكن أن تعمل المعالجة على نفسها. يمكن حل العديد من الألغاز الفلسفية التقليدية المتعلقة بمعرفة النفس (وخداع النفس) باستخدام تلك الفكرة المألوفة للذكاء الاصطناعي.

لكن ما مدى الاضطلاع بـ «معرفة النفس»؟ ينكر بعض الفلاسفة حقيقة النفس، ولكن المفكرين المتأثرين بالذكاء الاصطناعي لا ينكرونها. إنهم يرونها نوعاً خاصاً من الأجهزة الافتراضية.

في رأيهم، النفس بنية حاسوبية دائمة تنظم أفعال العامل وتهذبها، لا سيما أفعاله التطوعية المدروسة بعناية. (على سبيل المثال، يصف مؤلف مشروع تعلم وكيل التوزيع الذكي النفس بأنها «السياق الدائم للتجربة التي تنظم التجارب وتثبتها عبر العديد من السياقات المحلية المختلفة»). إنها ليست موجودة في الطفل الحديث الولادة، ولكنها بناء يستمر بطول العمر، كما أنها طيعة إلى حد ما إلى القولة الذاتية المتعمدة. يتيح تعدد أبعادها تنوعاً كبيراً؛ مما يولد عاملاً فردياً يمكن تمييزه، وخاصية شخصية مميزة.

السبب المحتمل في ذلك هو أن نظرية عقل العامل (التي تفسر سلوك الآخرين مبدئياً) تُطبّق من دون تفكير على أفكار الشخص الفردية وأفعاله. إنها تُعطي معنى منطقياً لأفعاله من حيث الدوافع والنوايا والأهداف التي لها الأولوية. وهذه بدورها تترتب حسب التفضيلات الفردية المستمرة والعلاقات الشخصية والقيم الأخلاقية/السياسية. تسمح هذه المعمارية الحاسوبية ببناء كل من الصورة الذاتية (التي تمثل الشخصية التي يعتقد الشخص أنه عليها) والصورة الذاتية المثالية (التي يحب الشخص أن يكون عليها)، وكذلك الأفعال والعواطف التي تستند إليها الفروق بين الصورتين.

يطلق دينيت (المتأثر بمينسكي كثيراً) على النفس «مركز ثقل السرد»؛ أي إنها بنية (جهاز ظاهري) تولّد أفعال المرء وتحاول أن تشرحها حينما تسرد حياة المرء الخاصة، لا سيما علاقات المرء مع الآخرين. بالطبع هذا يترك مجالاً لخداع النفس واللامنظورية المتعددة الجوانب للنفس.

وبالمثل، يصف دوجلاس هوفشتادتر النفس بأنها أنماط مجردة ذاتية المرجعية تنشأ من أساس لا معنى له للنشاط العصبي ثم تعود إليه في بعض الأحيان. هذه الأنماط

(الأجهزة الافتراضية) ليست أنماطاً سطحية للشخص. على النقيض من ذلك، حتى توجد النفس فلا يلزم سوى تجسيد ذلك النمط.

باختصار، اتخاذ القرار بارتقاء أنظمة الذكاء الاصطناعي العام إلى مستوى ذكاء الإنسان — بما في ذلك الأخلاقيات والحرية والنفس — سيكون خطوة كبيرة لها تأثيرات عملية خطيرة. ومن المرجح أن يكون هؤلاء الذين ينكر حدسهم الفكرة بالجملة لأنها خاطئة من الأساس مصيبين. للأسف، لا يمكن دعم حدسهم بحجج فلسفية غير مثيرة للجدل. لا يوجد إجماع على تلك المسائل؛ ومن ثم لا توجد إجابات سهلة.

العقل والحياة

كل العقول التي نعرفها موجودة في الكائنات الحية. ويعتقد كثيرون — بمن فيهم علماء السبرانية (انظر الفصلين الأول والخامس) — أنه لا بد من وجود العقل. أي إنهم يفترضون أن العقل يقتضي وجود الحياة مسبقاً.

يصرّح الفلاسفة المهنيون في بعض الأحيان بهذا، ولكن نادراً ما يحاجّون عنه. قال بوتنام — على سبيل المثال — مما لا شك فيه أنه إذا لم يكن الروبوت حياً فلن يكون واعياً. لكنه لم يقدم أسباباً علمية، بل اعتمد على «القواعد الدلالية للغة». حتى القلة الذين دافعوا عن هذا الافتراض مدة طويلة لم يقدموا دليلاً لا يدع مجالاً للشك، ومنهم الفيلسوف البيئي هانس يونس، ومؤخراً عالم الفيزياء كارل فريستون عبر «مبدأ الطاقة الحرة» السيبراني العام.

لكن لنفترض أن هذا الاعتقاد السائد صحيح. وإذا كان صحيحاً، فلن يتحقق الذكاء الحقيقي في أنظمة الذكاء الاصطناعي إلا إذا دبّت فيها حياة حقيقية. عندئذٍ، لا بد أن نسأل ما إذا كانت «الحياة الاصطناعية القوية» (الحياة في الفضاء السيبراني) ممكنة أم لا.

لا يوجد تعريف للحياة مقبول لدى الجميع. ولكن عادةً ما تُذكر تسع سمات للحياة، وهي: التنظيم الذاتي، والاستقلالية، والنشوء، والنمو، والتكيف، والاستجابة، والتكاثر، والتطور، والأيض. يمكن فهم أول ثمان سمات من حيث معالجة المعلومات؛ ومن ثمّ يمكن تجسيدها في الذكاء الاصطناعي/الحياة الاصطناعية. على سبيل المثال، تحقق التنظيم الذاتي بطرق متنوعة، وقد فهم على نطاق واسع، كما أنه يتضمن كل السمات الأخرى (انظر الفصلين الرابع والخامس).

لكن هل هو ذكاء حقاً؟

لكن الأيض مختلف. يمكن نمذجة الأيض بالكمبيوتر، ولكن لا يمكن تمثيله به. فلا الروبوتات الذاتية التجميع ولا الحياة الاصطناعية الافتراضية (على الشاشة) يمكنها عمل التمثيل الغذائي في الواقع. الأيض هو استخدام المواد الكيميائية الحيوية وعمليات تبادل الطاقة من أجل تكوين الكائن الحي والحفاظ عليه. ومن ثم فهي عملية مادية لا يمكن اختزالها. يشير المدافعون عن الحياة الاصطناعية القوية إلى أن أجهزة الكمبيوتر تستخدم الطاقة، وأن بعض الروبوتات لديها مخازن طاقة فردية وتحتاج إلى التجديد الدوري. ولكن هذا بعيدٌ كل البعد عن الاستخدام المرن للدورات الكيميائية الحيوية المتشابكة التي تبني النسيج الجسدي للكائن الحي.

لذا، إذا كان الأيض ضرورياً من أجل الحياة، فإن الحياة الاصطناعية القوية يستحيل تحقيقها. وإذا كانت الحياة ضرورية من أجل العقل، فإن الحياة الاصطناعية القوية يستحيل تحقيقها كذلك. لا يهم مدى روعة أداء بعض أنظمة الذكاء الاصطناعي العام في المستقبل، لكن لن يكون لها ذكاء حقاً.

الانقسام الفلسفي الكبير

يعتبر الفلاسفة التحليليون ومعهم الباحثون في الذكاء الاصطناعي أن احتمالية بعض الجوانب النفسية العلمية أمرٌ مسلمٌ به. وفي الحقيقة، افترض هذا الموقف بين صفحات الكتاب، بما في ذلك هذا الفصل.

لكن علماء الظواهر يرفضون هذا الافتراض. يقولون إن كل المفاهيم العلمية ناشئة من وعي له مغزى؛ ومن ثم لا يمكن استخدامها لشرح هذا الافتراض. وقبل وفاة بوتمان عام ٢٠١٦، هو نفسه قبل هذا الموقف. حتى إنهم يزعمون أنه من غير المنطقي افتراض وجود عالم حقيقي قائم بشكل مستقل عن الفكر البشري، وربما يكتشف العلم خصائصه الموضوعية.

إذن، عدم الإجماع على رأي بشأن طبيعة العقل/الذكاء أعمق مما أشرت إليه في هذا الكتاب.

لا يوجد حجة دامغة تدحض رأي علماء الظواهر، ولا تؤيدها كذلك. يرجع السبب إلى عدم وجود أرض مشتركة يمكن الانطلاق منها. فكل فريق يدافع عن نفسه وينتقد الفريق الآخر، ولكنهم يستخدمون حججاً لا يتفق الطرفان على مصطلحاتها الأساسية. تقدّم الفلسفة التحليلية والفلسفة الظاهرية تفسيرات مختلفة اختلافاً جوهرياً حتى للمفاهيم

الذكاء الاصطناعي

الأساسية، مثل «السبب» و«الحقيقة». (قدّم عالم الذكاء الاصطناعي براين كانتويل سميث ميتافيزيقا طموحة لكلّ من الحوسبة والقصدية والكائنات التي تهدف إلى احترام رؤى كلا الفريقين، ولكن مع الأسف لم تكن حجته المثيرة للجدل مقنعة).
لم يُحلّ هذا النزاع، وقد لا يُحلّ. يرى البعض أن رأي علماء الظواهر جليّ الصواب. ويراه آخرون أنه جليّ السخف. ولهذا السبب لا أحد يعلم على وجه اليقين ما إذا كان الذكاء الاصطناعي العام سيصبح ذكيًا بحق أم لا.

الفصل السابع

التفرد

أثيرت ضجة حول الذكاء الاصطناعي منذ ظهوره. كذلك أثارت التوقعات الحماسية للغاية من (بعض) العاملين في الذكاء الاصطناعي حماسة الصحفيين ومقدمي البرامج الثقافية، كما أثارت رعبهم في بعض الأحيان. واليوم، المثال الأساسي هو التفرد؛ ويعني وقتاً معيناً من الزمن تصبح الآلة فيه أذكى من الإنسان.

في البداية، قيل إن الذكاء الاصطناعي سيبلغ مستوى ذكاء الإنسان. (ويُفترض ضمناً أنه سيكون ذكاءً حقيقياً؛ انظر الفصل السادس). لم يمضِ وقتٌ طويل حتى تحوّل مصطلح الذكاء الاصطناعي العام إلى الذكاء الاصطناعي الخارق. ستتمتع الأنظمة بذكاء يمكنها من نسخ نفسها؛ ومن ثم يتفوق عددها على عدد البشر، وأن تحسن نفسها؛ ومن ثم يتفوق ذكاؤها على ذكاء الإنسان. عندئذٍ، ستحل أهم المسائل وتتخذ أهم القرارات بواسطة أجهزة الكمبيوتر.

هذه الفكرة مثيرة للجدل إلى حد كبير. يختلف الناس بشأن هل يمكن أن يحدث هذا، وهل سيحدث أم لا، ومتى قد يحدث، وهل سيكون شيئاً جيداً أم سيئاً.

يقول المؤمنون بالتفرد إن التقدم المحرز في مجال الذكاء الاصطناعي يجعل التفرد أمراً لا مفر منه. رُحِبَ البعض بهذا. إنهم يتوقعون حل مشكلات البشرية. الحروب والأمراض والمجاعات والضجر وحتى القتل ... ستنتهي كل تلك المشكلات. يتوقع آخرون نهاية البشرية، أو على أي حال ستنتهي الحياة المدنية التي نعرفها. أثار ستيفين هوكينج (ومعه ستيوارت راسل، المؤلف المشارك في كتاب رائد عن الذكاء الاصطناعي) موجةً عالمية في مايو ٢٠١٤ حينما قال: إن تجاهل تهديدات الذكاء الاصطناعي «قد تكون أسوأ خطأ لنا على الإطلاق».

وعلى النقيض من ذلك، لا يتوقع المشككون في التفرد أن يحدث التفرد، وبالتأكيد ليس في المستقبل القريب. يقولون إن الذكاء الاصطناعي يثير كثيرًا من القلاقل. لكنهم لا يرون أي تهديد وجودي.

رُسل التفرد

في الآونة الأخيرة، شاعت فكرة التحول من الذكاء الاصطناعي العام إلى الذكاء الاصطناعي الخارق في وسائل الإعلام، ولكنها بدأت في منتصف القرن العشرين. المبادرون الأساسيون هم جاك جود (زميل ألان تورينج الذي فك الشفرة معه في بلتشلي بارك) وفيرنور فينج وراي كورزويل. (توقع تورينج نفسه أن تتولى الآلة زمام الأمور، ولكنه لم يصرّح).

في عام ١٩٦٥، تنبأ جود بآلة فائقة الذكاء، وتوقع أن تتفوق تلك الآلة إلى حد بعيد في كل الأنشطة العقلية لأي إنسان مهما بلغ ذكاؤه. وبما أنها يمكن أن تصمم آلات ذات ذكاء أفضل، فلا ريب في أنها [ستؤدي إلى] ثورة في الذكاء. حينذاك، كان جود متفائلًا، ولكن على وجل: «الآلة الفائقة الذكاء هي آخر اختراع يحتاج الإنسان إلى صنعه، ولكن بشرط أن تكون تلك الآلة طيعة بالقدر الذي يجعلها نخبرنا كيف نبقيها تحت تصرفنا.» لكن بعد ذلك، احتجّ بقوله إن الآلات الفائقة الذكاء ستدمرنا.

بعد ربع قرن من الزمان، أشاع فينج مصطلح «التفرد» (وأول من أطلقه في هذا السياق هو جون فون نيومان عام ١٩٥٨). تنبأ بـ «التفرد التكنولوجي القادم» الذي ستنهال عنده كل التنبؤات (مقارنةً بأفق الحدث في الثقوب السوداء).

سلم بأن التفرد نفسه يمكن توقعه، بل إنه لا مفر منه في الحقيقة. ولكن قد يكون له تبعات كثيرة (غير معروفة)، ومنها احتمالية تدمير الحضارة وحتى الجنس البشري. إننا نتجه نحو «التخلص من كل القواعد السابقة، وربما يحدث ذلك في غمضة عين، إنه هروب فائق السرعة يتخطى أي أمل في السيطرة.» حتى إن أدركت كل الحكومات الخطر وحاولت أن تمنعه، فلن تستطيع.

تصدى كروزويل لتشاؤم فينج ومن بعده جود (في نهاية المطاف). فهو لم يقدم فآل خير فحسب، بل أعطى مواعيد محدّدة.

في كتابه الذي يحمل عنوان «التفرد قريب»، يقول إن الذكاء الاصطناعي العام سيتحقق بحلول عام ٢٠٣٠؛ وبحلول عام ٢٠٤٥، سيكون الذكاء الاصطناعي الخارق (بالإضافة إلى تكنولوجيا النانو والتكنولوجيا الحيوية) قد تغلب على الحروب والأمراض

والفقر وأسباب الموت. كذلك سيؤدي إلى «ثورة في الفن والعلوم وأشكال المعارف الأخرى التي ستعطي معنى للحياة حقاً». وبحلول منتصف القرن أيضاً، سنعيش في واقع افتراضي أكثر ثراءً وإرضاءً من العالم الحقيقي. وفي رأي كورزويل، «التفرد» فريد حقاً، و«قريب» تعني أنه قريب بالفعل.

(يخفُّ هذا التفاؤل المفرط في بعض الأحيان. وسرد كورزويل العديد من المخاطر التي تهدد الوجود والناجمة بنسبة كبيرة من التكنولوجيا الحيوية التي تعمل بمساعدة الذكاء الاصطناعي. وأما بشأن الذكاء الاصطناعي نفسه، فإنه يقول: «الذكاء بطبيعته يستحيل السيطرة عليه ... ولا جدوى اليوم من وضع استراتيجيات تضمن تماماً أن يتحلّى الذكاء الاصطناعي في المستقبل بالأخلاق والقيم»).

استندت حجة كورزويل على قانون مور؛ إذ لاحظ مؤسس شركة إنتل جوردون مور أن قوة الكمبيوتر المتاحة مقابل دولار واحد تتضاعف كل عام. (ستتغلب قوانين الفيزياء على قانون مور في النهاية، ولكن ليس في المستقبل القريب). وكما أشار كورزويل، أي زيادة فائقة السرعة لا تأتي على حسب المتوقع تماماً. ويقول هنا، إن هذا ينطوي على أن الذكاء الاصطناعي يتقدم بسرعة لا يمكن تخيلها. ومثل فينج، يُصرُّ على أن التوقعات المبنية على التجارب السابقة لا قيمة لها تقريباً.

تنبؤات متنافسة

على الرغم من التصريح بعدم قيمة التنبؤات بشأن ما بعد التفرد، غالباً ما تُطرح هذه التنبؤات. توجد مجموعة من الأمثلة المحيرة للعقل في المؤلفات، ولكننا نذكر قليلاً منها في هذا الكتاب.

ينقسم المؤمنون بالتفرد إلى معسكرين؛ المتفائلين (من أتباع فينج) والمتشائمين (من أتباع كورزويل). يتفق معظمهم في أن التحول من الذكاء الاصطناعي العام إلى الذكاء الاصطناعي الخارق سيقع قبل نهاية هذا القرن. ولكنهم يختلفون فقط في العواقب الوخيمة التي سيؤدي إليها الذكاء الاصطناعي الخارق.

على سبيل المثال، يتوقع البعض أن تفعل الروبوتات الشريرة كل ما يحطم آمال البشر وحياتهم (وهذا النموذج شائع في أفلام الخيال العلمي وأفلام هوليوود). وعلى وجه التحديد، أنكرت فكرة أنه يمكننا منع هذا الخراب إذا لزم الأمر. يقال إن أنظمة الذكاء الاصطناعي الخارق ستتحلّى بالذكاء الذي يمكّنها من الحيلولة دون ذلك.

يعترض آخرون ويقولون إن أنظمة الذكاء الاصطناعي الخارق لن يكون لها نوايا خبيثة، «ولكنها ستكون خطيرة على أي حال». إننا لن نغرس كراهية البشر في تلك الأنظمة، ولا يوجد سبب يدعوهم إلى اكتساب تلك الكراهية من تلقاء نفسها. بل إنهم لن يبالوا بنا، مثلما نفعل نحن تجاه معظم الأنواع من غير البشر. لكن عدم المبالاة من جانبها قد يكون زلة من جانبنا إن تعارضت مصالحنا مع الأهداف التي رُكبت فيها. الإنسان العاقل مثل طائر الدودو. في التجربة التي تخيلها نيك بوستروم والمقتبسة على نطاق واسع، ينوي نظام الذكاء الاصطناعي الخارق عمل مشابه ورق من شأنها أن تبحث عن الذرات في أجسام البشر كي تصنعها.

مرة أخرى، فكر في استراتيجية عامة تُقترح أحياناً للحماية من تهديدات التفرد، ألا وهي الاحتواء. وهنا، يُمنع الذكاء الاصطناعي الخارق من أن يتصرف في العالم من تلقاء نفسه على الرغم من أن بإمكانه التعرف على العالم من تلقاء نفسه. إنه لا يُستخدم إلا للإجابة عن أسئلتنا (وهو ما يُسميه بوستروم «أوراكل»). لكن العالم يوجد به شبكة الإنترنت، وربما تُحدث أنظمة الذكاء الاصطناعي الخارق تغييرات بطريقة غير مباشرة عن طريق المساهمة بالمحتوى على الإنترنت، مثل الوقائع والأكاذيب وفيروسات الكمبيوتر. شكل آخر من أشكال التشاؤم من التفرد، وهو أن الآلات ستضطرننا إلى القيام بأعمالها القذرة، حتى إن كان هذا يتعارض مع صالح البشرية. وهذه الفكرة تنكر فكرة أنه يمكننا «احتواء» أنظمة الذكاء الاصطناعي الخارق عن طريق الحيلولة دون اتصالها بالعالم. يقال إن الآلة الفائقة الذكاء يمكن أن تستخدم الرشوة أو التهديدات كي تقنع أحد الأشخاص القلائل التي تتصل بهم في بعض الأحيان لأداء أشياء يتعذر عليها أدائها بنفسها.

يفترض هذا القلق على وجه التحديد أن يتعلم الذكاء الاصطناعي الخارق مقداراً كبيراً عن الجانب النفسي لدى الإنسان، لدرجة أن يعرف أنواع الرشوة أو التهديد الذي يؤثر، وربما أيضاً يتعرف على الأشخاص الذين يصلح معهم نوع معين من الإقناع. ورداً بعدم قبول صحة هذا الافتراض، فإن الرشاوى المالية المغرية أو التهديدات بالقتل تكاد تصلح مع أي شخص؛ ومن ثم لن يحتاج الذكاء الاصطناعي الخارق إلى رؤية نفسية تُباري الرؤية النفسية التي لدى هنري جيمس. وكذلك لن تحتاج إلى فهم ما تعنيه كلمات الإقناع والرشوة والتهديد بالمعنى الذي اصطلح عليه البشر. لن يحتاج النظام سوى معرفة أن إدخال نصوص معينة بمعالجة اللغة الطبيعية والتحدث بها إلى الإنسان يمكن أن تؤثر في سلوكه بطرق يسهل التنبؤ بها.

بعض التنبؤات المتفائلة أصعب بكثير. ربما أكثر ما يلفت الانتباه إليها توقعات كورزويل عن العيش في الحياة الافتراضية وغياب فكرة أن يخطف الموت أحبابنا. قد يستمر موت الجسد على الرغم من تطاول الأعمار (بفضل علوم الأحياء التي تستخدم الذكاء الاصطناعي الخارق). ولكن يمكن التخلص من فكرة الموت بتحميل الشخصيات وذكريات الأفراد على أجهزة الكمبيوتر.

برز هذا الافتراض المثير للإشكاليات الفلسفية بأن الشخص يمكن أن يوجد في مادة سيليكون أو في مادة بروتين عصبية (انظر الفصل السادس) في كتاب ألفه كورزويل عام ٢٠٠٥ بعنوان: «عندما يتفوق البشر على علم الأحياء». عبّر كورزويل عن رؤيته «التفردية» — ويطلق عليها أيضًا الإنسانية العنصرية أو بعد الإنسانية — لعالم يتكون جزئيًا أو حتى كليًا من أشخاص ليسوا من الكائنات البشرية.

يُزعم أن هؤلاء «الأشخاص المسيرين آليًا» في حقبة ما بعد الإنسانية، سيوضع لهم غرسات حاسوبية تتصل مباشرة بالدماع والأطراف الصناعية والأعضاء الحسية أو أي من ذلك. سيزول العمى والصمم؛ لأن الإشارات البصرية والسمعية ستُفسّر من خلال حاسة اللمس. على الأقل، سيتعزز الإدراك العقلي (وكذلك الحالة المزاجية) بفضل أدوية مصممة خصيصًا.

الإصدارات الأولى من هذه التكنولوجيات المساعدة موجودة بالفعل. وإذا كانت تنتشر بسرعة على حد قول كورزويل، فسيُغيّر مفهومنا عن الإنسانية تغييرًا عميقًا. وبدلاً من النظر إلى الأطراف الصناعية باعتبارها أدوات مساعدة لجسم الإنسان، سترى على أنها أجزاء (متحولة) من جسم الإنسان. ستُدرج العقاقير المخمرة للعقل إلى جنب المواد الطبيعية مثل الدوبامين ضمن حسابات «الدماع». وسيُعتبر الذكاء الخارق أو القوة أو الجمال لدى الأفراد المعدّلين وراثيًا سمات «طبيعية». كذلك ستُنقّض الآراء السياسية بشأن المساواة والديمقراطية. وأيضًا ربما تتطور أنواع فرعية (أو أنواع) من أسلاف الإنسان الثرية بدرجة تتيح لها استغلال هذه الاحتمالات.

باختصار، يتوقع أن يحل التطور التقني محل التطور البيولوجي. يرى كورزويل أن التفرد هو «ذروة اندماج تفكيرنا البيولوجي ووجودنا مع التكنولوجيا؛ مما يؤدي إلى عالم [لا] يكون فيه تمييز ... بين الإنسان والآلة أو بين الواقع المادي والواقع الافتراضي». (ربما يُغفر لك شعورك بالحاجة إلى نفس عميق).

الإنسانية العنصرية مثال متطرف عن المدى الذي يمكن أن يغيّر به الذكاء الاصطناعي الأفكار بشأن طبيعة الإنسان. إن الفلسفة الأقل تطرفًا التي تستوعب التكنولوجيا في

مفهوم العقل هي «العقل الممتد»، حيث ترى العقل متغلغلًا في جميع أرجاء العالم بحيث يشمل العمليات المعرفية التي تعتمد عليه. وعلى الرغم من الانتشار الواسع لفكرة العقل الممتد، فإن فكرة الإنسانية العنصرية ليست كذلك. وقد تحمّس جمعٌ من الفلاسفة ومقدمي البرامج الثقافية والفنانين لدعم تلك الفكرة. وعلى الرغم من ذلك، لا يقبلها كل المؤمنون بالتفرد.

الدفاع عن الشك

في رأيي، التشكيك بشأن التفرد في محله. ورد في مناقشة تصوير العقل بالآلة الافتراضية بالفصل السادس أنه لا توجد عقبات أمام مبدأ بلوغ الذكاء الاصطناعي حد الذكاء البشري (وربما يستثنى من ذلك الوعي المُدرَك بالحواس). والسؤال هنا يدور حول ما إذا كان يمكن تنفيذ هذا عمليًا.

بالإضافة إلى اللامعقولية البديهية للعديد من التنبؤات بشأن مرحلة ما بعد التفرد وفلسفة الإنسانية العنصرية التي تكاد تصل إلى حد العبث (في رأيي)، فإن المشككين في التفرد لديهم حجج أخرى.

الذكاء الاصطناعي ليس مبشرًا بالقدر الذي يعتقدونه الكثيرون. ورد في الفصول من الثاني حتى الخامس أشياء عديدة لا يمكن للذكاء الاصطناعي أدائها في الوقت الحالي. والعديد من تلك الأشياء يتطلب إحساسًا بشريًا بالملاءمة (ويفترض ضمنيًا إكمال الويب الدلالي؛ انظر الفصل الثاني). إضافة إلى ذلك، ركّز الذكاء الاصطناعي على العقلانية الفكرية وتجاهل الذكاء الاجتماعي/العاطفي؛ ناهيك عن الحكمة. فالذكاء الاصطناعي العام الذي يتفاعل مع عالمنا تفاعلًا تامًا سيحتاج إلى تلك القدرات أيضًا. أضف إلى ذلك الثراء الهائل للعقول البشرية والحاجة إلى نظريات نفسية/حاسوبية وحيية عن طريقة عملها، في حين أن آفاق الذكاء الاصطناعي العام على مستوى الإنسان تبدو قاتمة. حتى وإن كان تنفيذه ممكنًا عمليًا، ثمة شكوك بشأن تحقيق التمويل اللازم. تخصص الحكومات حاليًا موارد هائلة في محاكاة الدماغ (انظر القسم التالي)، لكن الأموال المطلوبة لبناء عقول بشرية اصطناعية ستكون أكبر بكثير.

بفضل قانون مور، فلا ريب من توقُّع المزيد من التقدم في مجال الذكاء الاصطناعي. لكن الزيادات في قوة أجهزة الكمبيوتر ووفرة البيانات (بفضل التخزين السحابي والمستشعرات التي تعمل على مدار الساعة وحتى إنترنت الأشياء) لا يضمن أن يبلغ الذكاء

الاصطناعي مستوى ذكاء الإنسان. تلك أخبار سيئة للمؤمنين بفكرة التفرد أن الذكاء الاصطناعي الخارق يحتاج إلى الذكاء الاصطناعي العام أولاً.

يتناسى المؤمنون بفكرة التفرد جوانب القصور التي يتَّسم بها الذكاء الاصطناعي في الوقت الراهن. إنهم ببساطة لا يبالون لأن لديهم ورقة رابحة، ألا وهي أن التقدم التكنولوجي السريع يعيد كتابة كل كتب القواعد. وهذا يخوِّلهم لطرح تنبؤات كما يشاءون. لكنهم يسلِّمون في بعض الأحيان بعدم واقعية التنبؤات التي ستحدث بحلول «نهاية القرن». وعلى الرغم من ذلك، يُصرُّون على أن كلمة «مستحيل» تعني وقتاً طويلاً. كلمة «مستحيل» تعني وقتاً طويلاً حقاً. ومن ثم، قد يخطئ المشكِّكون في التفرد، ومنهم أنا. فهم ليس لديهم حجة دامغة، خاصة إذا سلِّموا باحتمالية الذكاء الاصطناعي العام من حيث المبدأ (كما أفعل أنا). وربما يقتنعون بأن التفرد سيحدث في النهاية حتى وإن طالت غيبته.

ومع ذلك، فإن الدراسة المتأنية للذكاء الاصطناعي الحديث تعطي سبباً وجيهاً لدعم فرضية المشكِّكين (أو رهانهم إذا كنت تفضل هذا المصطلح) بدلاً من التكهّنات الجامحة للمؤمنين بالتفرد.

المحاكاة الكاملة للدماغ

يتوقع المؤمنون بالتفرد أن يحدث تقدم تكنولوجي فائق السرعة في الذكاء الاصطناعي والتكنولوجيا الحيوية وتكنولوجيا النانو، وفي التعاون فيما بين هذه المجالات. وفي الحقيقة، هذا التقدم واقع الآن. تُستخدم تحليلات البيانات الضخمة لإحراز تقدم في الهندسة الوراثية وتطوير الأدوية، والعديد من المشروعات العلمية الأخرى (وبذلك يُسوِّغ لرأي آدا لافليس؛ انظر الفصل الأول). وبالمثل، يُمزج الذكاء الاصطناعي وعلم الأعصاب في المحاكاة الكاملة للدماغ.

تهدف المحاكاة الكاملة للدماغ إلى محاكاة دماغ حقيقية عن طريق محاكاة مكوناته الفردية (الخلايا العصبية)، بالإضافة إلى وصلاتها وقدراتها الخاصة بمعالجة المعلومات. يؤمل أن يكون للمعرفة العلمية المكتسبة في العديد من المجالات، بما في ذلك علاج الأمراض العقلية بدءاً من ألزهايمر وصولاً إلى الفصام.

ستتطلب هذه الهندسة العكسية الحوسبة العصبية، توفر نماذج للعمليات الخلوية الفرعية مثل مرور الأيونات عبر غشاء الخلية (انظر الفصل الرابع).

تعتمد الحوسبة العصبية على المعرفة بتشريح الأنواع المختلفة من الخلايا العصبية والعلم بوظائفها. لكن المحاكاة الكاملة للدماغ ستطلب أدلة تفصيلية عن وصلات عصبية محددة وعن وظائفها، بما في ذلك التوقيت. سيتطلب قدرٌ كبير من هذه العملية فحصاً محسناً للدماغ بالإضافة إلى مسابير عصبية مصغرة لا تتوقف عن رصد الخلايا العصبية الفردية.

في الوقت الراهن، تُنفَّذ العديد من مشروعات المحاكاة الكاملة للدماغ، وكثيراً ما يقارنها رعاتها بمشروع الجينوم البشري أو السباق إلى القمر. على سبيل المثال، أعلن الاتحاد الأوروبي عام ٢٠١٣ عن مشروع الدماغ البشري بتكلفة تبلغ مليار جنيه إسترليني. وفي وقت لاحق من العام نفسه، أعلن الرئيس الأمريكي باراك أوباما عن مبادرة أبحاث الدماغ عبر النهوض بالتقنيات العصبية المبتكرة، وقد استمر المشروع ١٠ سنوات، وموّلته الحكومة الأمريكية بمبلغ ثلاثة مليارات دولار أمريكي (بالإضافة إلى مبلغ كبير من الأموال الخاصة). يهدف المشروع في المقام الأول إلى تصميم خريطة ديناميكية للاتصالية في أدمغة الفئران، ثم محاكاة حالة الإنسان.

الحكومة هي الجهة التي موّلت المحاولات السابقة للمحاكاة الجزئية للدماغ. عام ٢٠٠٥، رعت سويسرا مشروع «الدماغ الأزرق»، وكان هدفه الأساسي محاكاة العمود القشري لدى الفأر، ولكن الهدف الطويل المدى هو إعداد نماذج للميون عمود قشري في القشرة الجديدة للإنسان. في عام ٢٠٠٨، قدّمت وكالة مشروعات الأبحاث المتطورة الدفاعية (داربا) نحو ٤٠ مليون دولار لمشروع «أنظمة الإلكترونيات العصبية والبلاستيكية القابلة للتطور»؛ وفي عام ٢٠١٤، وبتقديم ٤٠ مليون دولار إضافية تقريباً، كان هذا المشروع يستخدم رقائق تحمل ٥,٤ مليارات ترانزستور، وكل ترانزستور يحمل مليون وحدة (خلية عصبية) و٢٥٦ مليون تشابك عصبي. كذلك تتعاون ألمانيا واليابان في استخدام «تكنولوجيا المحاكاة العصبية» لتطوير الكمبيوتر كيه؛ وبحلول عام ٢٠١٢، كان الأمر يستغرق ٤٠ دقيقة لمحاكاة ثانية واحدة من واحد بالمائة من نشاط الدماغ الحقيقي، وكان يتضمن ١,٧٣ مليار خلية عصبية و١٠,٤ تريليونات تشابك عصبي.

ونظراً إلى التكلفة الباهظة، نادراً ما تُجرى المحاكاة الكاملة للدماغ على الثدييات. ولكن تُنفَّذ محاولات لا حصر لها حول العالم لتخطيط الأدمغة الأصغر بكثير (في جامعتي، هم يركزون على نحل العسل). وربما تقدم هذه التجارب رؤى علمية عصبية يمكن أن تساعد في المحاكاة الكاملة للدماغ على مستوى الإنسان.

وانطلاقاً من التقدم المتحقق بالفعل في مكونات الأجهزة (مثل رقاقات أنظمة الإلكترونيات العصبية والبلاستيكية القابلة للتطوير) بالإضافة إلى قانون مور، فإنه يمكن تصديق تنبؤ كورزويل بأن أجهزة الكمبيوتر المطابقة لقدرة المعالجة الأولية في دماغ الإنسان ستوجد في عشرينيات القرن الحادي والعشرين. ولكن اعتقاده بأنها ستتطابق مع ذكاء الإنسان بحلول ٢٠٣٠ مسألة أخرى.

ما يهمنا في هذا المقام هي الأجهزة الافتراضية (انظر الفصلين الأول والسادس). بعض الأجهزة الافتراضية لا يمكن تنفيذها إلا بمكونات أجهزة بالغة القوة. لذلك قد تكون رقائق الكمبيوتر ذات عدد الترانزستورات الضخم ضرورية. لكن ما عمليات الحوسبة التي ستنفذها؟ بعبارة أخرى، ما الأجهزة الافتراضية التي ستنفذ بتلك المكونات؟ لمضاهاة ذكاء الإنسان (أو حتى الفأر)، لا بد أن تكون هذه الأجهزة قوية معلوماتياً بطرق لم يفهمها علماء النفس الحاسوبيون فهمًا تامًا حتى الآن.

لنفترض أن كل خلية عصبية في دماغ الإنسان سيجري تخطيطها في النهاية، ولكني لا أظن أن هذا سيحدث. هذا في حد ذاته لن يخبرنا بما سيحدث. (لا تحتوي الدودة الخيطية الصغيرة «الربداء الرشيق» إلا على ٣٠٢ خلية عصبية، والاتصالات بين تلك الخلايا معروفة بدقة. ولكن لا يمكننا حتى تحديد هل هذه التشابكات العصبية محفزة أم مثبطة).

بالنسبة إلى القشرة البصرية، لدينا الآن بالفعل تخطيط تفصيلي بين التشريح العصبي والوظائف النفسية. ليس الأمر كذلك بالنسبة إلى القشرة الدماغية الجديدة بوجه عام. وعلى وجه التحديد، لا نعرف كثيرًا عما تفعله القشرة الأمامية؛ بمعنى ماهية الأجهزة الافتراضية التي تنفذ فيها. هذا السؤال ليس بارزًا في المحاكاة الكاملة للدماغ الواسعة النطاق. على سبيل المثال، اعتمد «مشروع الدماغ البشري» نهجًا تصاعديًا؛ انظر إلى التشريح والكيمياء الحيوية وحاول محاكتهما. هُملت الأسئلة المتنوعة بشأن الوظائف النفسية التي قد يدعمها الدماغ (ولم يشترك سوى قلة من علماء علم الأعصاب المعرفي). حتى إن تحققت النمذجة التشريحية بالكامل، ورُصدت الرسائل الكيميائية بعناية، فلن تتم الإجابة عن تلك الأسئلة.

قد تطلب تلك الإجابات مجموعة كبيرة من المفاهيم الحاسوبية. إضافة إلى ذلك، من الموضوعات الأساسية في تلك المسألة المعمارية الحاسوبية للدماغ (أو العقل والدماغ) ككل. رأينا في الفصل الثالث أن تخطيط العمل لدى المخلوقات التي لها عدة أهداف تتطلب

آليات جدولة معقدة، مثل الآليات التي توفرها العاطفة. وأشارت المناقشة بشأن نموذج تعلّم وكيل التوزيع الذكي في الفصل السادس إلى التعقيد البالغ في المعالجة القشرية. حتى النشاط الدنيوي المتمثل في تناول الطعام بالسكين والشوكة يتطلب دمج العديد من الأجهزة الافتراضية، وبعض تلك الأجهزة يتعامل مع أشياء مادية (العضلات والأصابع والأوتان وأنواع مختلفة من المستشعرات)، والبعض الآخر مع النوايا والخطط والتوقعات والرغبات والتقاليد الاجتماعية والتفضيلات. لفهم مدى احتمالية ذلك النشاط، لا نحتاج بيانات علمية عصبية عن الدماغ فحسب، بل نحتاج إلى نظريات حاسوبية تفصيلية عن العمليات النفسية المشتركة في النشاط.

باختصار، إن اعتبرنا المحاكاة الكاملة للدماغ طريقة لفهم الذكاء البشري، فمن المرجح أن تفشل المحاكاة الكاملة التصاعدية للدماغ. قد تعلمنا الكثير عن الدماغ. وقد تساعد علماء الذكاء الاصطناعي على تطوير تطبيقات عملية أكثر. ولكن فكرة أن المحاكاة الكاملة للدماغ ستشرح الذكاء البشري بحلول منتصف القرن ما هي إلا وهم.

ما ينبغي أن نقلق بشأنه

إذا كان المشككون في التفرد على صواب، ولن يتحقق التفرد، فلا يعني ذلك أنه لا يوجد ما نقلق إزاءه. فالذكاء الاصطناعي يثير المخاوف بالفعل. وسيثير التقدم المستقبلي المزيد من القلق؛ ومن ثمّ القلق بشأن السلامة في الذكاء الاصطناعي على المدى البعيد في محله تمامًا. أضف إلى ذلك أننا بحاجة إلى أن ننتبه إلى تأثيراته على المدى القصير أيضًا.

بعض المخاوف عادية إلى حد كبير. على سبيل المثال، يمكن استخدام أي تكنولوجيا في الخير أو الشر. سيستخدم الخبيثون أي أدوات متاحة — وربما يمولون تطوير أدوات جديدة — لارتكاب أفعال مشينة. (مشروع «سي واي سي» على سبيل المثال قد يكون مفيدًا لمرتكبي الأفعال الشريرة؛ والمطورون بالفعل يفكرون في طرق لتقييد الوصول إلى النظام الكامل عند إصداره، انظر الفصل الثاني). لذا يجب أن نتحرى أقصى درجات الحذر فيما نخترع.

وكما أشار ستيفارت راسل، هذا يعني ما هو أكبر من مجرد الانتباه إلى أهدافنا. إذا كانت هناك ١٠ معلومات ذات صلة بمسألة ما، وإحصائيًا تحسين تعلم الآلة لا يراعي أكثر من ست معلومات (انظر الفصل الثاني)، فيمكن دفع المعلومات الأربعة المتبقية إلى حدود متطرفة. وهنا، يجب أن نفطن أيضًا إلى أنواع البيانات التي تُستخدم.

هذا القلق العام معني بمسألة الإطار (انظر الفصل الثاني). ومثل الصياد في القصة الخيالية، حينما تمنى أن يعود ابنه الجندي إلى الوطن ولكنه عاد إليه محمولاً على الأعناق، ربما نندهش كثيراً من أنظمة الذكاء الاصطناعي القوية التي تفتقر إلى فهمنا لماهية الصلة بالموضوع.

على سبيل المثال، عندما أوصى نظام الإنذار المبكر للحرب الباردة (في الخامس من أكتوبر ١٩٦٠) بضربة دفاعية ضد الاتحاد السوفيتي، لم تُتجنب الكارثة إلا لما أحسَّ المشغلون بالأهمية السياسية والإنسانية على حد سواء. قالوا إن السوفيت في الأمم المتحدة لم يحدثوا صخباً لا سيما في الآونة الأخيرة، وقد خشوا أن تحدث عواقب وخيمة جرّاء هجوم نووي. ومن ثمَّ كسروا البروتوكولات وتجاهلوا التحذير الآلي. كادت أن تقع العديد من الحوادث النووية الأخرى، وبعضها لم يمر عليه وقت طويل. عادةً لا يُمنع التصعيد إلا بالمنطق السليم لدى الناس.

إضافة إلى ذلك، وقوع أخطاء من الإنسان محتمل على الدوام. تلك الأخطاء تكون مفهومة في بعض الأحيان. (تفاقت حالة الطوارئ في مشروع جزيرة الثلاثة أميال حينما أوقف البشر جهاز الكمبيوتر، ولكن الظروف المادية التي كانوا يواجهونها لم تكن طبيعية إلى حد كبير). ولكن ربما تكون غير متوقعة إلى حد الاندهاش. صار التحذير من الحرب الباردة المذكور في الفقرة السابقة واقعاً لأنَّ أحدًا نسي السنوات الكبيسة عند برمجة التقويم؛ ومن ثمَّ كان القمر في المكان الخطأ. وكل هذا يضيف سبباً إلى ضرورة اختبار برامج الذكاء الاصطناعي وإثبات موثوقيتها (إن أمكن) قبل الاستخدام.

المخاوف الأخرى محدّدة أكثر. لا بد أن بعضها يُزعجنا اليوم.

البطالة بسبب التكنولوجيا واحدة من أهم التهديدات. فقد اختفت العديد من الوظائف اليدوية والكتابية المنخفضة المستوى. ولكن سيعقب ذلك وظائف أخرى (على الرغم من أن الوظائف اليدوية التي تتطلب البراعة والقدرة على التكيف لن تختفي). والآن، يمكن القيام بمعظم أعمال الرفع والجلب والحمل في المخازن باستخدام الروبوتات. وتطوير سيارات بدون سائق ستعني فقدان الناس لوظائفهم.

حتى الوظائف الإدارية الوسطى عُرضة للاختفاء هي الأخرى. يستخدم العديد من الاحترافيين الآن أنظمة الذكاء الاصطناعي باعتبارها أدوات مساعدة. لن يمر وقتٌ طويل حتى يستولي الذكاء الاصطناعي على جزء كبير من الوظائف (في القانون والمحاسبة، على سبيل المثال) التي تتضمن عمليات بحث طويلة في اللوائح والأعمال السابقة. كذلك سرعان

ما ستتأثر المهن ذات درجة الطلب العالية، ومنها المهن الطبية والعلمية. ستقل المهارات في الوظائف حتى وإن لم تُفقد. سيتهور التدريب المهني؛ كيف سيتعلم الفتيان إصدار أحكام منطقية؟

في حين أن بعض الوظائف القانونية ستزيد عن الحاجة، فإن المحامين سيستفيدون من الذكاء الاصطناعي؛ لأن مجموعة من الفخاخ القانونية تصطفُ منتظرة. إذا حدث خطأ ما، فمن سيكون المسئول؛ هل المبرمج أم تاجر الجملة أم بائع التجزئة أم المستخدم؟ وهل يمكن مقاضاة أحد المهنيين على عدم استخدامه نظام الذكاء الاصطناعي؟ إذا ثبت أن النظام يمتاز بدرجة موثوقية عالية (سواء رياضياً أو تجريبياً)، فالأرجح أن يحدث هذا التقاضي.

لا شك أنه ستظهر أنواع وظائف جديدة. لكن المشكوك فيه هو ما إذا كانت هذه الوظائف ستعادل الوظائف القديمة من حيث الأعداد وإمكانية الوصول إلى الوسائل التعليمية وقدرة المعيلين (مثلما حدث بعد الثورة الصناعية) أم لا. فنحن نتظرنا تحديات اجتماعية وسياسية خطيرة.

لن يتعرض قطاع الخدمات إلى تهديد كبير. ولكن حتى هذه الوظائف معرضة للخطر. وفي أفضل الحالات، ستُغتتم الفرص بحماسة للمضاعفة والارتقاء بالأنشطة بين الأشخاص التي لا تحظى بالتقدير الكافي في الوقت الحالي. ولكن هذا ليس مضموناً.

على سبيل المثال، يُفتح التعليم لعوامل المساعدة الفردية والذكاء الاصطناعي القائم على الإنترنت أو أحدهما، ومن ذلك الدورات التدريبية الضخمة المفتوحة على الإنترنت التي توفر محاضرات يقدمها أكاديميون بارزون؛ مما يحط من قدر المهارات الوظيفية لدى العديد من المعلمين في الفصول. يتوفر الأطباء النفسيون على الكمبيوتر بالفعل، وبتكلفة أقل من الأطباء في العيادات. (بعض هؤلاء الأطباء مفيدون إلى حد كبير، في التعرف على الاكتئاب، على سبيل المثال). وعلى الرغم من ذلك، فتلك الخدمات لا تخضع لأي لوائح قانونية البتة. وقد رأينا في الفصل الثالث أن التغير الديموجرافي يشجع على البحث في مجال ربما يصبح مربحاً، وهو «تقديم خدمات الرعاية» بالذكاء الاصطناعي لكبار السن وكذلك «المربيات الروبوتية».

وبصرف النظر عن أي تأثيرات على البطالة، فإن استخدام أنظمة الذكاء الاصطناعي الخالي من العاطفة في هذه السياقات الإنسانية الجوهرية محفوف بالمخاطر من الناحية العملية ومريب من الناحية الأخلاقية. صُمم العديد من «أجهزة الكمبيوتر المرافقة» كي

يستخدمها كبار السن أو المعاقين ممن لهم تواصل ضئيل مع عدد قليل من البشر الذين يتعاملون معهم. ليس الهدف من هذه الأجهزة أن تكون مجرد مصادر للمساعدة والترفيه فحسب، بل لتكون أيضًا وسائل للمحادثة والمؤانسة والارتياح العاطفي. حتى إذا زادت سعادة الإنسان المعرض للخطر بتلك التكنولوجيا (مثل سعادة مستخدمي الروبوت بارو)، فإنه يتخلّى عن كرامتهم الإنسانية رويدًا رويدًا. (الاختلافات الثقافية مهمة في هذا السياق؛ تختلف المواقف تجاه الروبوتات اختلافًا كبيرًا بين اليابان والغرب على سبيل المثال).

قد يستمتع المستخدمون من كبار السن بمناقشة ذكرياتهم الشخصية مع مرافق ذي ذكاء اصطناعي. لكن هل هذه مناقشة حقًا؟ قد تكون تذكيرًا مرغوبًا، حيث إنها تحيي نوبات مطمئنة من الحنين إلى الماضي. ومع ذلك، يمكن تقديم هذه الميزة دون إغراء المستخدم بوهم التعاطف. وفي كثير من الأحيان، ما يريده الشخص قبل كل شيء هو الاعتراف بشجاعته و/أو معاناته، حتى في مواقف الاستشارات المشحونة بالعاطفة. ولكن ينشأ هذا من الفهم المشترك للحالة الإنسانية. إننا نولي الأفراد اهتمامًا أقل من المتوقع عندما لا نقدم لهم سوى محاكاة سطحية للتعاطف.

حتى إذا كان المستخدم يعاني الخرف إلى حد ما، فمن المرجح أن تكون «نظريته» عن عامل الذكاء الاصطناعي أثرى من نموذج العامل البشري. إذن، ما النتيجة إن أخفق العامل في الاستجابة حسب المتوقع وحسب الحاجة عندما يستغرق الشخص في ذكريات فقد آلمته (كأن يكون فقد ولدًا)؟ لن تجدي تعبيرات التعاطف في المحادثة مع الرفيق نفعا، وربما تضر أكثر مما تنفع. وفي هذه الأثناء، قد يتفاهم كرب الشخص من دون أن تتوافر له السلوى على الفور.

هناك قلق آخر، وهو هل الرفيق عليه أن يصمت في بعض الأحيان أم يقول كذبة بيضاء. قول الحقيقة الصارمة (والصمت المفاجئ) قد يُغضب المستخدم. لكن اللباقة تتطلب تطويرًا هائلًا في معالجة اللغات الطبيعية بالإضافة إلى نموذج دقيق لعلم النفس البشري.

بالنسبة إلى المربيات الروبوتات (وتجاهل مشكلات السلامة)، فإن الإفراط في استخدام أنظمة الذكاء الاصطناعي مع الأطفال والرضع قد يؤدي إلى تشوّه نموهم الاجتماعي واللغوي أو أحدهما.

لا يُصوّر شركاء الجنس الاصطناعيون في الأفلام فحسب (مثل فيلم «هي» (هير)). بل يُسوَّق لها بالفعل. بعضها يستطيع التعرف على الكلام، وكذلك يستطيع إصدار كلام

أو حركات إغراء. إنها تزيد من تأثيرات الإنترنت التي تزيد من العنف الجنسي لدى البشر في الوقت الحالي (وتعزز التشويه الجنسي للمرأة). كتب الكثيرون (ومنهم علماء في الذكاء الاصطناعي) عن اللقاءات الجنسية مع الروبوتات بعبارات تكشف عن مفهوم ضحل للغاية عن الحب، ويكاد يخلط بينه وبين الشهوة والهوس الجنسي ومجرد الألفة الباعثة على الراحة. ولكن لا يُحتمل أن تكون هذه الملاحظات التحذيرية فعالة. بناءً على الريح الذي تُدِرُّه صناعة الإباحية بوجه عام، فهناك أمل ضئيل في منع «التطورات» المستقبلية في دُمى الجنس ذات الذكاء الاصطناعي.

الخصوصية مشكلة أخرى معقدة. تزداد حدة الجدل في المسألة، حيث إن بحث الذكاء الاصطناعي القوي وتعلم الذكاء الاصطناعي متروكان للبيانات التي تُجمع من الوسائط الشخصية وأجهزة الاستشعار في المنازل أو المخازن. (حصلت جوجل على براءة اختراع في دب روبوتي، حيث وضعت الكاميرا في عينيه، والميكروفون في أذنيه، ومكبرات الصوت في فمه. سيكون قادرًا على التواصل مع الوالدين والأطفال، وكذلك مع جامعي البيانات غير المرئيين أيضًا بصورة عشوائية).

الأمن السيبراني مشكلة منذ أمد بعيد. وكلما توغلَّ الذكاء الاصطناعي في عالمنا (وكثيرًا ما يحدث ذلك بطرق غير شفافة)، زادت أهميته. من وسائل الدفاع ضد الاستحواذ على الذكاء الاصطناعي الخارق هو إيجاد طرق لكتابة خوارزميات لا يمكن اختراقها/تغييرها (وهو هدف الذكاء الاصطناعي الصديق؛ انظر القسم التالي).

التطبيقات العسكرية أيضًا تثير المخاوف. الروبوتات الكاسحات الألغام مرحَّبٌ بها بشدة. لكن ماذا عن الجنود الروبوتيين أو الأسلحة الروبوتية؟ الطائرات من دون طيار الحالية يشغلها الإنسان، لكنها قد تزيد المعاناة من خلال توسيع المسافة البشرية (وليس الجغرافية فقط) بين المشغل والهدف. ولا بد أن نرجو ألا يُسمح للطائرات بدون طيار أن تقرر مَنْ/ما الهدف المنشود. حتى الثقة بها في التعرف على هدف (جرى اختياره بشريًا) تثير قضايا مقلقة من الناحية الأخلاقية.

الإجراءات الوقائية من تلك المخاوف

هذه المخاوف ليست جديدة، على الرغم من أنه لم يلاحظها من العاملين في الذكاء الاصطناعي سوى قلة حتى الآن.

تدارس العديد من رواد الذكاء الاصطناعي التأثيرات الاجتماعية في اجتماع عُقد في ليك كومو عام ١٩٧٢، لكن رفض جون مكارثي أن ينضم إليهم، وقال إن الأمر لا يزال باكرًا على التفكير. بعد بضع سنوات، نشر عالم الكمبيوتر جوزيف فايزنباوم كتابًا بعنوان «من الحكم إلى الحساب»، يندد فيه بـ «الفحش» الناتج عن الخلط بين الاثنين، ولكن نبذه مجتمع الذكاء الاصطناعي مستهزئًا به.

كانت هناك بعض التوقعات بالفعل. على سبيل المثال، في أول كتاب يُسلط الضوء على الذكاء الاصطناعي (من مؤلفاتي، ونُشر عام ١٩٧٧)، تناول الفصل الأخير موضوع «الأهمية الاجتماعية». تأسست مؤسسة «متخصصو الكمبيوتر من أجل المسؤولية الاجتماعية» عام ١٩٨٣ (وجزء من تأسيسها راجع إلى جهود تيري فينوجراد مطور برنامج «شردلو»؛ انظر الفصل الثالث). لكن هذا فُعل في المقام الأول للتحذير من عدم موثوقية تكنولوجيا حرب النجوم، حتى إن عالم الكمبيوتر ديفيد بارناس خاطب مجلس الشيوخ الأمريكي حول هذا الموضوع. وعندما انحسرت مخاوف الحرب الباردة، بدا معظم محترفي الذكاء الاصطناعي وكأنه قلَّ شغفهم بمجال عملهم. ولم يستمر في التركيز على القضايا الاجتماعية/الأخلاقية على مدار سنوات سوى قلة منهم، مثل نيويل شاركي من جامعة شيفيلد (عالم روبوتات ترأس اللجنة الدولية للحد من أسلحة الروبوت)، بالإضافة إلى بعض فلاسفة الذكاء الاصطناعي، مثل ويندل والاش من جامعة ييل، وبلاي ويتبي من جامعة ساسكس.

والآن، نظرًا إلى الممارسات والمستقبل الواعد للذكاء الاصطناعي، أصبحت هذه الهواجس ملحّة أكثر. ومن ثَم أصبحت التداعيات الاجتماعية داخل المجال (وإلى حد ما خارجه) تنال مزيدًا من الاهتمام.

بعض الاستجابات المهمة لا علاقة لها بالتفرد. على سبيل المثال، تدعو الأمم المتحدة ومنظمة رصد حقوق الإنسان منذ فترة إلى معاهدة (لم تُوقّع حتى الآن) تحظر الأسلحة المستقلة بالكامل، مثل الطائرات من دون طيار التي تختار الهدف. وفي الآونة الأخيرة، راجعت بعض الهيئات الراسخة أولويات البحث ومدونة قواعد السلوك لديها أو أحدهما. ولكن الحديث عن التفرد جر المزيد من المساهمين إلى النقاش.

يقول كثيرون من المؤمنين بالتفرد والمشككين فيه على حد سواء إنه حتى وإن كانت احتمالية تحقيق التفرد ضئيلة للغاية، فإن التبعات المحتملة بالغة الخطر، وينبغي أن نتخذ التدابير الاحترازية بدءًا من الآن. وعلى الرغم من زعم فينج بأنه لا يمكن فعل شيء إزاء التهديدات الوجودية، أقيمت العديد من المؤسسات للحماية منها.

ومن هذه المؤسسات في المملكة المتحدة، مركز دراسات المخاطر الوجودية بكامبريدج ومعهد مستقبل الإنسانية بأكسفورد؛ وفي الولايات المتحدة، معهد مستقبل الحياة في بوسطن ومعهد أبحاث ذكاء الآلة في بيركلي.

تتلقى هذه المؤسسات تمويلات ضخمة من الأسخياء على الذكاء الاصطناعي. على سبيل المثال، شارك جان تالين في تطوير برنامج سكايب، وشارك في تأسيس مركز دراسات المخاطر الوجودية ومعهد مستقبل الحياة. تحاول هاتان المؤسستان أن تنبها صناعات السياسات وغيرهم من المؤثرين في الرأي العام إلى تلك المخاطر، هذا إلى جانب التواصل مع المتخصصين في الذكاء الاصطناعي.

نظم رئيس الجمعية الأمريكية للذكاء الاصطناعي (إريك هورويتز) لجنة مصغرة عام ٢٠٠٩ لمناقشة التدابير الاحترازية التي يمكن اتخاذها لتوجيه — أو حتى تأجيل — العمل في الذكاء الاصطناعي المثير للمشكلات الاجتماعية. ووضح أن هذه اللجنة عقدت في أسيلومار بكاليفورنيا، حيث اتفق متخصصون في علم الوراثة قبل بضع سنوات على إيقاف أبحاث معينة في علم الوراثة. ولكن بصفتي عضواً في المجموعة، ترسّخ لديّ انطباع أنه ليس كل المشاركين مهتمين بمستقبل الذكاء الاصطناعي. لم يحظَ التقرير المطمئن بانتشار مكثف في وسائل الإعلام.

عقد كلٌّ من معهد مستقبل الحياة ومركز دراسات المخاطر الوجودية اجتماعاً ذا دوافع مماثلة ولكنه أكبر (بموجب قواعد دار تشاتام وبدون وجود صحفيين) في بورتوريكو في يناير ٢٠١٥. شارك منظم الاجتماع ماكس تيجمارك في التوقيع مع راسل وهوكينج على الرسالة التحذيرية قبل ستة أشهر. لا عجب إذن من أن تكون الأجواء ملحة أكثر بكثير مما كانت عليه في أسيلومار. فقد أدت إلى توفير تمويل على الفور (من مليونير الإنترنت إيلون ماسك) للأبحاث بشأن سلامة الذكاء الاصطناعي والذكاء الاصطناعي الأخلاقي — بالإضافة إلى خطاب تحذيري مفتوح وقّع عليه آلاف العاملين في الذكاء الاصطناعي، وقد تناقلته وسائل الإعلام على نطاق واسع.

بعد فترة وجيزة، خطاب مفتوح آخر صاغه توم ميتشل والعديد من رواد الباحثين، حيث حذروا من تطوير الأسلحة المستقلة التي تحدد الأهداف وتتشابك معها من دون تدخل البشر. يأمل الموقعون في «منع انطلاق سباق تصنيع أسلحة الذكاء الاصطناعي». قدّم الخطاب في المؤتمر الدولي للذكاء الاصطناعي المنعقد في يوليو ٢٠١٥، وقد وقع عليه

ما يقرب من ٣ آلاف عالم و١٧ ألفاً من المتخصصين في مجالات ذات صلة، وقد كان له صدًى واسع في وسائل الإعلام.

أدى الاجتماع الذي عُقد في بورتوريكو إلى صياغة خطاب مفتوح (في يونيو ٢٠١٥)، وقد صاغه عالِم الاقتصاد من معهد ماساتشوستس للتكنولوجيا إريك برينجولفسون وأندي مكافي. كان هذا الخطاب موجَّهًا إلى صناع السياسات وأصحاب المشروعات ورجال الأعمال وكذلك علماء الاقتصاد. وتحذيرًا من الآثار الاقتصادية الخطيرة المحتملة جرَّاء الذكاء الاصطناعي، فقد اقترح الخطاب بعض التوصيات السياسية العامة التي ربما تخفف من وطأة المخاطر، على الرغم من أنها لن تمنعها.

شهد شهر يناير ٢٠١٧ دعوة أخرى لاجتماع بشأن الذكاء الاصطناعي المفيد. ومن جديد، نظم تيجمارك الاجتماع في مدينة أسيلومار الشهيرة.

تهدف هذه الجهود المبذولة من جانب مجتمع الذكاء الاصطناعي إلى إقناع الممولين الحكوميين عبر الأطلسي بأهمية القضايا الاجتماعية/الأخلاقية. ومؤخرًا، صرَّح كلٌّ من وزارة الدفاع الأمريكية ومؤسسة العلوم الوطنية عن رغبتهما في تمويل تلك الأبحاث. ولكن هذا الدعم ليس جديدًا بالكامل؛ فاهتمام الحكومات يتنامى منذ بضع سنوات.

على سبيل المثال، رعت مجالس الأبحاث في المملكة المتحدة أحد معارض «روبوتيكس ريتريت» المتعددة التخصصات عام ٢٠١٠، وشاركت في صياغة مدونة قواعد السلوك الخاصة بعلماء الروبوتات. جرى الاتفاق على خمسة «مبادئ»، تناول اثنان منها المخاوف التي ناقشناها مسبقًا، وهما: «(١) يجب ألا تُصمم الروبوتات كي تكون أسلحة، إلا لأسباب تتعلق بالأمن الوطني، (٤) الروبوتات مصنوعات يدوية؛ ومن ثم لا ينبغي استخدام وهم المشاعر والنية لاستغلال المستخدمين الضعفاء».

حمل مبدآن آخران البشر المسؤولية الأخلاقية المباشرة إذا أفادا بأن: «(٢) البشر — وليس الروبوتات — هم المسؤولون ... (٥) يجب أن تتوفر إمكانية معرفة المسئول [قانونًا] عن أي روبوت». امتنعت المجموعة عن محاولة تحديث «القوانين الثلاثة للروبوتات» التي وضعها إسحاق عظيموف (يجب ألا يؤذي الروبوت الإنسان، ويجب أن يطيع أوامر الإنسان، ويحمي حياته ما لم يتعارض هذا مع القانون الأول). وتؤكد المجموعة على ضرورة أن يتبع المصمم/المطور البشري أي «قوانين» يُنص عليها هنا.

في مايو ٢٠١٤، أشادت وسائل الإعلام بمبادرة أكاديمية مؤلَّتها البحرية الأمريكية (بمبلغ ٧,٥ ملايين دولار أمريكي لمدة خمس سنوات). وقد اشترك في هذا المشروع خمس

جامعات (ييل وبراون وتافتس وجورج تاون ومعهد رينسيلار)، ويهدف إلى تطوير «الكفاءة الأخلاقية» في الروبوتات. كذلك يضم المشروع علماء في علم النفس الاجتماعي والمعرفي وفي الفلسفة الأخلاقية، وكذلك مبرمجين ومهندسين في الذكاء الاصطناعي. لا تحاول هذه المجموعة المتعددة التخصصات أن تقدم قائمة بخوارزميات أخلاقية (مقارنة بقوانين عظيموف)، ولا أن ترتب أخلاقاً فوقية حسب الأولوية (مثل مذهب المنفعة)، ولا حتى أن تضع مجموعة من القيم الأخلاقية غير المتنافسة. بل ترجو أن يطور نظام حاسوبي قادر على التفكير المنطقي الأخلاقي (والمناقشة الأخلاقية) في العالم الواقعي. فالروبوتات المستقلة ستتخذ قرارات تداولية في بعض الأحيان، ولن تتبع فقط التعليمات (غير أنها تتفاعل تفاعلاً غير مرّن مع الإشارات الكائنة؛ انظر الفصل الخامس). إذا شارك الروبوت في عملية بحث وإنقاذ على سبيل المثال، فمن الذي ينبغي أن يخليه من المكان/ينقذه أولاً؟ أو إذا كان يقدم رفقة اجتماعية، فمتى ينبغي ألا يخبر مستخدمه الحقيقة إذا اقتضت الضرورة ذلك؟

سيدمج النظام المقترح الإدراك والعمل الحركي ومعالجة اللغات الطبيعية والتفكير المنطقي (سواء الاستنتاجي أو التماثلي) والعاطفة. تتضمن العاطفة التفكير العاطفي (الذي يمكن أن يشير إلى أحداث مهمة، وكذلك يُجدول الأهداف المتضاربة؛ انظر الفصل الثالث)، والعروض الروبوتية التي تُظهر «الاحتجاج والضيق» ما يمكن أن يؤثر في القرارات الأخلاقية التي يتخذها المتفاعل مع النظام، والتعرف على العواطف لدى البشر المحيطين بالروبوت. كذلك صرّح الإعلان الرسمي أن الروبوت يمكن أن يتخطى حتى الحد العادي (أي البشري) في الكفاءة الأخلاقية.

بناءً على العقبات التي تقف أمام الذكاء الاصطناعي العام الواردة في الفصلين الثاني والثالث، بالإضافة إلى الصعوبات المتعلقة بالأخلاقيات على وجه التحديد (انظر الفصل السادس)، بمقدور المرء أن يشك في إمكانية تحقيق هذه المهمة. لكن قد يكون المشروع مفيداً على الرغم من ذلك. بالتفكير في مشكلات العالم الواقعي (مثل المثلين المختلفين تماماً الواردين فيما سبق)، فقد ينبّهنا إلى العديد من مخاطر استخدام الذكاء الاصطناعي في المواقف ذات الإشكالية الأخلاقية.

بجانب هذه الجهود المؤسسية، يتزايد عدد فرادى علماء الذكاء الاصطناعي الذين يسعون إلى تحقيق ما يطلق عليه أليازر يودكوفسكي «الذكاء الاصطناعي الصديق». إنه ذكاء اصطناعي له تأثيرات إيجابية على البشر، فهو آمن ومفيد. يشتمل النظام على

خوارزميات واضحة وموثوقة ودقيقة، وتفشل من دون ضرر إن حدث وفشلت. يجب أن تكون الخوارزميات شفافة، ويمكن التنبؤ بها، وليست عرضة لأن يتلاعب بها المخترقون؛ وإذا أمكن إثبات موثوقيتها بالمنطق أو الرياضيات بدلاً من الاختبار التجريبي، فهذا أفضل بكثير.

تبرّع ماسك بمبلغ ستة ملايين دولار أمريكي في اجتماع بورتوريكو؛ مما أدى إلى إعلان «دعوة للاقتراحات» غير مسبوق من معهد مستقبل الحياة (وبعد ستة أشهر، بلغ عدد المشروعات الممولة ٣٧ مشروعاً). باستهداف الخبراء في السياسة العامة والقانون والأخلاقيات والاقتصاد والتعليم والتواصل وكذلك الذكاء الاصطناعي، فقد طالب الاجتماع بالآتي: «أن تهدف المشروعات البحثية إلى رفع مستوى الفائدة الاجتماعية المستقبلية المرجوة من الذكاء الاصطناعي، وفي الوقت نفسه تجنب المخاطر المحتملة، والاقتصار على الأبحاث التي لا تركز صراحةً على الهدف القياسي، وهو رفع قدرات الذكاء الاصطناعي، بل تركز على رفع مستوى قوة الذكاء الاصطناعي وفائدته أو أحدهما». ربما حدث هذا النداء المرحب به للذكاء الاصطناعي الصديق على أي حال. لكن أثر التفرد كان واضحاً: «ستعطى الأولوية للأبحاث التي تهدف إلى الحفاظ على دقة الذكاء الاصطناعي وفائدته، حتى وإن وصل الأمر إلى إبطال قدر كبير من القدرات الحالية».

باختصار، الرؤى شبه المروعة لمستقبل الذكاء الاصطناعي وهمية. ولكن جزءاً من أسبابها راجع إلى تنبؤ مجتمع الذكاء الاصطناعي وصناع السياسات وكذلك العامة إلى بعض المخاطر الحقيقية بالفعل. ينبغي التحرك الآن؛ إذ قد تأخرنا كثيراً.

قراءات إضافية

Boden, M. A. (2006), *Mind as Machine: A History of Cognitive Science*, 2 vols. (Oxford: Oxford University Press).

With the exception of deep learning and the Singularity, every topic mentioned in this *Very Short Introduction* is discussed at greater length in *Mind as Machine*.

Russell, S., and Norvig, P. (2013), *Artificial Intelligence: A Modern Approach*, 3rd edn. (London: Pearson).

This is the leading textbook on AI.

Frankish, K., and Ramsey, W., eds. (2014), *Cambridge Handbook of Artificial Intelligence* (Cambridge: Cambridge University Press).

Describes the various areas of AI, less technically than Russell and Norvig (2013).

Whitby, B. (1996), *Reflections on Artificial Intelligence: The Social, Legal, and Moral Dimensions* (Oxford: Intellect Books).

A discussion of aspects of AI that are too often ignored.

Husbands, P., Holland, O., and Wheeler, M. W., eds. (2008), *The Mechanical Mind in History* (Cambridge, MA: MIT Press).

The fourteen chapters (and five interviews with AI/A-Life pioneers) describe early work in AI and cybernetics.

Clark, A. J. (1989), *Microcognition: Philosophy, Cognitive Science, and Parallel Distributed Processing* (Cambridge, MA: MIT Press).

An account of the differences between symbolic AI and neural networks.

Today's neural networks are much more complex than those discussed here, but the main points of comparison still stand.

Minsky, M. L. (2006), *The Emotion Machine: Commonsense Thinking, Artificial Intelligence, and the Future of the Human Mind* (New York: Simon & Schuster).

This book, by one of the founders of AI, uses AI ideas to illuminate the nature of everyday thought and experience.

Hansell, G. R., and Grassie, W., eds. (2011), *H +/-: Transhumanism and Its Critics* (Philadelphia: Metanexus).

Statements and critiques of the transhumanist philosophy supported, and the transhumanist future predicted, by some AI visionaries.

Dreyfus, H. L. (1992), *What Computers Still Can't Do: A Critique of Artificial Reason*, 2nd edn. (New York: Harper and Row).

The classic attack, based in Heideggerian philosophy, of the very idea of AI. (Know your enemies!)

المراجع

NB: ‘MasM’, in the chapter references that follow, identifies the most relevant parts of Boden, *Mind as Machine*.

(For the analytical table of contents of MasM, see the ‘Key Publications’ section of my website: “www.ruskin.tv/margaretboden”.)

الفصل الأول: ما الذكاء الاصطناعي؟

MasM chaps. 1.i.a, 3.ii–v, 4, 6.iii–iv, 10–11.

The quotations from Ada Lovelace are from: Lovelace, A. A. (1843), ‘Notes by the Translator’. Reprinted in R. A. Hyman (ed.) (1989), *Science and Reform: Selected Works of Charles Babbage* (Cambridge: Cambridge University Press), 267–311.

Blake, D. V., and Uttley, A.M. (eds.) (1959), *The Mechanization of Thought Processes*, vol. 1 (London: Her Majesty’s Stationery Office).

This contains several important early papers, including descriptions of Pandemonium and Perceptrons and a discussion of AI and common sense.

McCulloch, W. S., and Pitts, W. H. (1943), ‘A Logical Calculus of the Ideas Immanent in Nervous Activity’, *Bulletin of Mathematical Biophysics*,

5: 115–33. Reprinted in S. Papert, ed. (1965), *Embodiments of Mind* (Cambridge, MA: MIT Press), 19–39.

Feigenbaum, E. A., and Feldman, J. A., eds. (1963), *Computers and Thought* (New York: McGraw–Hill).

An influential collection of early papers on AI.

الفصل الثاني: كأن الذكاء العام هو الكأس المقدسة

MasM, sections. 6.iii, 7.iv, and chaps. 10, 11, 13.

Boukhtouta, A. et al. (2005), *Description and Analysis of Military Planning Systems* (Quebec: Canadian Defence and Development Technical Report).

Shows how AI planning has advanced since the early days.

Mnih, V., and D. Hassabis et al. (nineteen authors) (2015), ‘Human–Level Control Through Deep Reinforcement Learning’, *Nature*, 518: 529–33.

This paper by the DeepMind team describes the Atari game–player.

Silver, D., and D. Hassabis et al. (seventeen authors) (2017), ‘Mastering the Game of Go Without Human Knowledge’, *Nature*, 550: 354–9.

This describes the latest version of DeepMind’s (2016), *AlphaGo* program (for the earlier version, see *Nature*, 529: 484–9).

The quotation from Allen Newell and Herbert Simon is from their (1972) book *Human Problem Solving* (Englewood–Cliffs, NJ: Prentice–Hall).

The quotation ‘new paradigms are needed’ is from LeCun, Y., Bengio, Y., and Hinton, G. E. (2015), ‘Deep Learning’, *Nature*, 521: 436–44.

Minsky, M. L. (1956), ‘Steps Toward Artificial Intelligence’. First published as an MIT technical report: *Heuristic Aspects of the Artificial Intelligence Problem*, and widely reprinted later.

Laird, J. E., Newell, A., and Rosenbloom, P. (1987), ‘Soar: An Architecture for General Intelligence’, *Artificial Intelligence*, 33: 1–64.

الفصل الثالث: اللغة، الإبداع، العاطفة

MasM, chaps. 7.ii, 9.x–xi, 13.iv, 7.i.d–f.

Baker, S. (2012), *Final Jeopardy: The Story of WATSON, the Computer That Will Transform Our World* (Boston: Mariner Books).

A readable, though uncritical, account of an interesting Big Data system.

Graves, A., Mohamed, A.–R., and Hinton, G. E. (2013), 'Speech Recognition with Deep Recurrent Neural Networks', *Proc. Int. Conf. on Acoustics, Speech, and Signal Processing*, 6645–49.

Collobert, R. et al. (six authors) (2011), 'Natural Language Processing (Almost) from Scratch', *Journal of Machine Learning Research*, 12: 2493–537.

The quotation describing syntax as superficial and redundant is from Wilks, Y. A., ed. (2005), *Language, Cohesion and Form: Margaret Masterman (1910–1986)* (Cambridge: Cambridge University Press), p. 266.

Bartlett, J., Reffin, J., Rumball, N., and Williamson, S. (2014), *Anti-Social Media* (London: DEMOS).

Boden, M. A. (2004), *The Creative Mind: Myths and Mechanisms*, 2nd edn. (London: Routledge).

Boden, M. A. (2010), *Creativity and Art: Three Roads to Surprise* (Oxford: Oxford University Press). This collection of twelve papers is largely about computer art.

Simon, H. A. (1967), 'Motivational and Emotional Controls of Cognition', *Psychological Review*, 74: 39–79.

Sloman, A. (2001), 'Beyond Shallow Models of Emotion', *Cognitive Processing: International Quarterly of Cognitive Science*, 2: 177–98.

Wright, I. P., and Sloman, A. (1997), *MINDER: An Implementation of a Protoemotional Architecture*, available at <ftp://ftp.cs.bham.ac.uk/pub/tech-reports/1997/CSRP-97-01.ps.gz>.

الفصل الرابع: الشبكات العصبية الاصطناعية

MasM chaps. 12, 14.

Rumelhart, D. E., and J. L. McClelland, eds. (1986), *Parallel Distributed Processing: Explorations in the Microstructure of Cognition*, vol. 1: *Foundations* (Cambridge, MA: MIT Press).

The whole volume is relevant; but the past-tense learner (written by Rumelhart and McClelland) is described on pp. 216–71.

Clark, A. (2016), *Surfing Uncertainty: Prediction, Action, and the Embodied Mind* (Oxford: Oxford University Press). A review of Bayesian approaches in cognitive science.

See also the paper by LeCun et al., and the two references from Demis Hassabis' team, cited under Chapter 2, earlier.

The two quotations about the network scandal are from Minsky, M. L., and Papert, S. A. (1988), *Perceptrons: An Introduction to Computational Geometry*, 2nd edn. (Cambridge, MA: MIT Press), viii–xv and 247–80.

Philippides, A., Husbands, P., Smith, T., and O'Shea, M. (2005), 'Flexible Couplings Diffusing Neuromodulators and Adaptive Robotics', *Artificial Life*, 11: 139–60.

A description of GasNets.

Cooper, R., Schwartz, M., Yule, P., and Shallice, T. (2005), 'The Simulation of Action Disorganization in Complex Activities of Daily Living', *Cognitive Neuropsychology*, 22: 959–1004.

Describes a computer model of Shallice's hybrid theory of action.

Dayan, P., and Abbott, L. F. (2001), *Theoretical Neuroscience: Computational and Mathematical Modelling of Neural Systems* (Cambridge, MA: MIT Press).

This volume does not discuss technological AI, but shows how ideas from AI are influencing the study of the brain.

الفصل الخامس: الروبوتات والحياة الاصطناعية

MasM chaps 4.v–viii and 15.

Beer, R. DS. (1990), *Intelligence as Adaptive Behavior: An Experiment in Computational Neuroethology* (Boston: Academic Press).

Webb, B. (1996), 'A Cricket Robot', *Scientific American*, 275(6): 94–9.

Brooks, R. A. (1991), 'Intelligence without Representation', *Artificial Intelligence*, 47: 139–59.

The seminal paper of situated robotics.

Kirsh, D. (1991), 'Today the Earwig, Tomorrow Man?', *Artificial Intelligence*, 47: 161–84.

A sceptical response to situated robotics.

Harvey, I., Husbands, P., and Cliff, D. (1994), 'Seeing the Light: Artificial Evolution, Real Vision', *From Animals to Animats 3* (Cambridge, MA: MIT Press), 392–401.

Describes the evolution of an orientation detector in a robot.

Bird, J., and Layzell, P. (2002), 'The Evolved Radio and its Implications for Modelling the Evolution of Novel Sensors', *Proceedings of Congress on Evolutionary Computation*, CEC-2002, 1836–41.

Turk, G. (1991), 'Generating Textures on Arbitrary Surfaces Using Reaction-Diffusion', *Computer Graphics*, 25: 289–98.

Goodwin, B. C. (1994), *How the Leopard Changed Its Spots: The Evolution of Complexity* (Princeton University Press).

Langton, C. G. (1989), 'Artificial Life', in C. G. Langton (ed.), *Artificial Life* (Redwood City: Addison-Wesley), 1–47. Revd. version in M. A. Boden, ed. (1996), *The Philosophy of Artificial Life* (Oxford: Oxford University Press), 39–94.

The paper that defined 'artificial life'.

الفصل السادس: لكن هل هو ذكاء حَقَّاً؟

MasM chaps. 7.i.g, 16.

Turing, A. M. (1950), 'Computing Machinery and Intelligence', *Mind*, 59: 433–60.

The quotations regarding 'the hard problem' are from Chalmers, D. J. (1995), 'Facing up to the Problem of Consciousness', *Journal of Consciousness Studies*, 2: 200–19.

The quotation by J. A. Fodor is from his (1992), 'The Big Idea: Can There Be a Science of Mind?', *Times Literary Supplement*, 3 July: 5–7.

Franklin, S. (2007), 'A Foundational Architecture for Artificial General Intelligence', in B. Goertzel and P. Wang (eds.), *Advances in Artificial General Intelligence: Concepts, Architectures, and Algorithms* (Amsterdam: IOS Press), 36–54.

Dennett, D. C. (1991), *Consciousness Explained* (London: Allen Lane). Slo-
man, A., and Chrisley, R. L. (2003), 'Virtual Machines and Conscious-
ness', in O. Holland (ed.), *Machine Consciousness* (Exeter Imprint Aca-
demic), *Journal of Consciousness Studies*, special issue, 10(4): 133–72.

Putnam, H. (1960), 'Minds and Machines', in S. Hook (ed.), *Dimensions of Mind: A Symposium* (New York: New York University Press), 148–79.

The quotation about Physical Symbol Systems is from Newell, A., and Simon, H. A. (1972), *Human Problem Solving* (Englewood-Cliffs, NJ: Prentice-Hall).

- Gallagher, S. (2014), 'Phenomenology and Embodied Cognition', in L. Shapiro (ed.), *The Routledge Handbook of Embodied Cognition* (London: Routledge), 9–18.
- Dennett, D. C. (1984), *Elbow Room: The Varieties of Free Will Worth Wanting* (Cambridge, MA: MIT Press).
- Millikan, R. G. (1984), *Language, Thought, and Other Biological Categories: New Foundations for Realism* (Cambridge, MA: MIT Press).
- An evolutionary theory of intentionality.

فصل السابع: التفرد

- Kurzweil, R. (2005), *The Singularity is Near: When Humans Transcend Biology* (London: Penguin).
- Kurzweil, R. (2008), *The Age of Spiritual Machines: When Computers Exceed Human Intelligence* (London: Penguin).
- Bostrom, N. (2005), 'A History of Transhumanist Thought', *Journal of Evolution and Technology*, 14(1): 1–25.
- Shanahan, M. (2015), *The Technological Singularity* (Cambridge, MA: MIT Press).
- Ford, M. (2015), *The Rise of the Robots: Technology and the Threat of Mass Unemployment* (London: Oneworld Publications).
- Chace, C. (2018), *Artificial Intelligence and the Two Singularities* (London: Chapman and Hall/CRC Press).
- Bostrom, N. (2014), *Superintelligence: Paths, Dangers, Strategies* (Oxford: Oxford University Press).
- Wallach, W. (2015), *A Dangerous Master: How to Keep Technology from Slipping Beyond Our Control* (Oxford: Oxford University Press).

- Brynjolfsson, E., and McAfee, A. (2014), *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies* (New York: W. W. Norton).
- Wilks, Y. A., ed. (2010), *Close Engagements with Artificial Companions: Key Social, Psychological, Ethical, and Design Issues* (Amsterdam: John Benjamins).
- Boden, M. A. et al. (fourteen authors) (2011), 'Principles of Robotics: Regulating Robots in the Real World', available on EPSRC website: <www.epsrc.ac.uk/research/ourportfolio/themes/>.

مصادر الصور

- (2-1) Monkey and bananas problem: how does the monkey get the bananas? (Reprinted from M. A. Boden, *Artificial Intelligence and Natural Man* (1977: 387)).
- (6-1) A global workspace in a distributed system. (Adapted from p. 88 of B. J. Baars, *A Cognitive Theory of Consciousness* (Cambridge: Cambridge University Press, 1988) with kind permission).
- (6-2) Similarities between GW terms and other widespread concepts. (Adapted from p. 44 of B. J. Baars, *A Cognitive Theory of Consciousness* (Cambridge: Cambridge University Press, 1988) with kind permission).

