



هل ستسيطر الآلات التي يمكنها التعلّم على العالم؟

Yann LeCun^{1,2*}

¹معهد كورنات، جامعة نيويورك، نيويورك، الولايات المتحدة

²ميتا، فريق أبحاث الذكاء الاصطناعي الأساسي (FAIR)، نيويورك، الولايات المتحدة

المراجعون الصغار

OISIN

العمر: 12



ZI-AN

العمر: 8



يُعدّ التعلّم جزءًا لا يتجزأ من حياتنا وحياة كل الحيوانات، ولكن هل تدرك مدى روعة قدرتنا على التعلّم؟ عندما نحاول إنشاء آلات يمكنها التعلّم، تواجهنا أسئلة عميقة حول طبيعة الذكاء وطريقة عمله. في هذه المقالة، سأخبرك حول الشبكات الاصطناعية الخاصة التي تُسمى بالشبكات العصبية وتحاكي الدماغ لإنتاج سلوك ذكي. والشبكات العصبية جزء لا يتجزأ من الذكاء الاصطناعي وتُستخدم على نطاق واسع في العديد من التطبيقات اليومية، بدءًا من التعرف على الوجه وحتى القيادة الذاتية. وأنا على يقين من أن الشبكات العصبية ستؤدي دورًا رئيسيًا في حياتنا بالمستقبل، حيث ستجعلها أكثر راحة وستعزز فهمنا للألغاز الكبرى مثل مفهوم الذكاء وطريقة عمل الدماغ.

فاز البروفيسور LeCun Yann بجائزة تورينج في عام 2018 مشاركة مع البروفيسور جيفري هينتون والبروفيسور يوشوا بنجيو لتحقيق إنجازات في المجال الهندسي والمفاهيمي جعلت الشبكات العصبية العميقة جزءًا أساسيًا من الحوسبة.

هل الآلات أذكى من البشر؟

هل أنت ذكي؟ أعتقد أن الكل سيقول: "نعم". وماذا عن كلبك أو قطتك؟ وماذا عن هاتفك الذكي؟ على الرغم من أن الذكاء أمر بديهي لنا، فهو ظاهرة معقدة [1] وما من تعريف واضح للذكاء، كما لا يوجد مقياس مطلق لمدى ذكاء نظام ما. الأمر الذي نقوم به غالبًا في **الذكاء الاصطناعي** هو محاولة استنساخ أو تجاوز قدرات البشر والحيوانات، ولا سيما القدرات العقلية مثل فهم اللغة والتفكير.

يتسم الذكاء الاصطناعي بالتغير الدائم، فكلما تم حلّ مسألة معينة بحيث تتجاوز الآلات أداء البشر في مهمة معينة، لا يكون ذكاءً اصطناعيًا بعد الآن. والكثير من الأشياء التي كانت في العادة جزءًا من الذكاء الاصطناعي لم تعد جزءًا منه، مثل استخدام النظام العالي لتحديد المواقع (جي بي إس) لمعرفة أفضل طريق أو أداء عمليات حسابية معقدة. ومن خلال تقدّم الذكاء الاصطناعي، تعلمنا أيضًا أن البشر جيدون أو أذكى نسبيًا في أشياء معينة، ولكن سيئون كثيرًا في أمور أخرى. على سبيل المثال، حتى وقت قريب، كان البشر يعتقدون أنهم لا يمكن هزيمتهم في الألعاب اللوحية مثل الشطرنج ولعبة غو. ولكننا نعلم الآن أن أنظمة الذكاء الاصطناعي المتقدمة يمكنها هزيمة أفضل اللاعبين من البشر، ويمكنك الاطلاع على **هذا البرنامج الوثائقي** للتعرف على أول نظام ذكاء اصطناعي (ألفا جو) يهزم بطلاً عالميًا من البشر في عام 2015. واليوم، هناك مهام أخرى يمكن لأنظمة الذكاء الاصطناعي القيام بها بشكل أفضل من البشر مثل الترجمة الفورية والتعرف على الصور. ولكن حتى أذكى الأنظمة تعتبر عالية التخصص؛ فهي جيدة في نطاق ضيق للغاية من المهام [2].

ويجب أن تكون الآلات الذكية حقًا قادرة على فعل ما تفعله كل الحيوانات، وهو التعلم [3]. يعني التعلم التكيف أو استخدام المعرفة المكتسبة في تجارب لتقديم أداء جيد في مهمة ضمن سيناريوهات جديدة. يتعلم البشر والحيوانات بشكل طبيعي، ولكن كيف تتعلم الآلات؟ هذا هو بحث **تعلم الآلة**، وفيه تطور أنظمة **وخوارزميات** تستقبل البيانات وتستخدمها لمواصلة تعديل سلوكها وتحسين أدائها. وتعلم الآلة هو الأساس للعديد من التطبيقات، مثل الترجمة الفورية لصفحات الويب إلى لغات مختلفة والتحكم الصوتي والتعرف على الوجه في الهواتف الذكية وسيارات القيادة الذاتية وتحليل البيانات الطبية. ولكن أفضل آلات التعلم لا تزال لا تستطيع التعلم مثل البشر. على سبيل المثال، يستغرق البشر حوالي 20 إلى 30 ساعة لتعلم القيادة بشكل صحيح، ولكن لا توجد حاليًا آلة يمكنها القيادة ذاتيًا مثل البشر، ولا حتى بعد آلاف الساعات. وكان أحد الإنجازات الكبيرة لتعلم الآلة معرفة كيفية محاكاة الطريقة التي يتعلم بها الدماغ.

الذكاء الاصطناعي (ARTIFICIAL INTELLIGENCE (AI))

القدرة على استنساخ قدرات
البشر والحيوانات، بل وتجاوزها
باستخدام الآلات.

تعلم الآلة (MACHINE LEARNING (ML))

تعليم الآلات كيفية التعلم أو
تحسين أدائها نتيجة تجربة
سابقة بأقل قدر ممكن من
التدخل البشري.

الخوارزمية (ALGORITHM)

مجموعة تعليمات توجه الآلات
في طريقة العمل.

كيف يمكننا التعلم؟

العصبونات (NEURONS)

خلايا الدماغ التي تعالج المعلومات باستخدام الإشارات الكهربائية.

المشبك العصبي (SYNAPSE)

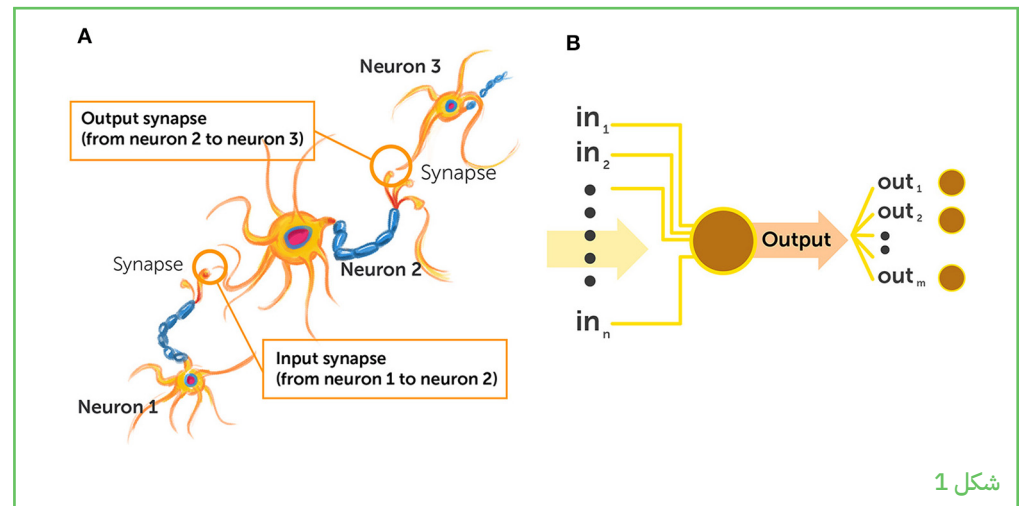
الوصلة التي تربط بين عصبونين. ويمكن أن تتفاوت قوة المشبك العصبي بمرور الوقت اعتمادًا على تكرار استخدامه لترميز الإشارة بين العصبونات.

تتألف الأدمغة من حوالي 100 مليار (مئة ألف مليون) عصبون وكل عصبون متصل بما يقرب من 10,000 عصبون آخر (معنى ذلك أنه لدينا نحو 1,000 تريليون اتصال في أدمغتنا!) وتتواصل **العصبونات** عبر الإشارات الكهربائية التي ترسلها وتستقبلها.

تعتمد قوة الإشارة الكهربائية التي يستقبلها كل عصبون من العصبونات الأخرى على **المشبك العصبي**، أي الوصلة التي تربطه بالعصبونات الأخرى (الشكل 1A؛ لمعرفة المزيد عن العصبونات والمشابك العصبية، راجع **هذه المقالة**). عند تعلّم شيء ما، تتغير قوة وموقع المشابك العصبية في دماغنا، فكّر في التعلّم كشبكة من الوحدات التي تستقبل مدخلات من وحدات أخرى وتؤدي عمليات حوسبة وتقوم بعمليات النقل والإخراج عبر وصلات تتغير قوتها بمرور الوقت. باستخدام هذه المبادئ، يمكننا إنشاء شبكة من **العصبونات الاصطناعية** التي يمكنها التعلم (الشكل 1B).

شكل 1

العصبونات الحيوية والاصطناعية: (A) العصبونات الحيوية في الدماغ التي تتغير قوة مشابكها العصبية وموقعها عندما نتعلم. (B) العصبونات الاصطناعية هي وحدات رقمية داخل الكمبيوتر تستقبل المدخلات وتنقل المخرجات إلى العصبونات الاصطناعية الأخرى.



شكل 1

العصبونات الاصطناعية (ARTIFICIAL NEURONS)

تمثيل للعصبونات داخل كمبيوتر. تستقبل العصبونات الاصطناعية المدخلات الرقمية وتقوم بعملية حسابية وتنقل مخرجات رقمية إلى العصبونات الاصطناعية الأخرى.

الشبكة العصبية (NEURAL NETWORK (NN))

شبكة من عناصر الحوسبة الاصطناعية تحاكي عمل العصبونات في الدماغ.

الشبكات العصبية الحيوية والاصطناعية

حدثت طفرة كبيرة في تعلم الآلة مع ظهور الشبكات العصبية [4]. **والشبكة العصبية** عبارة عن شبكة من العصبونات الاصطناعية المتصلة التي تتم محاكاتها في كمبيوتر وتنظيمها في طبقات (لمعرفة المزيد عن الشبكات العصبية، يمكنك قراءة **هذه المقالة** حول الشبكات العصبية الاصطناعية أو **هذه المقالة** حول التشابه بين الخلايا العصبية والذكاء الاصطناعي). والطبقة الأولى (طبقة المدخلات) (الشكل 2)، تستقبل البيانات التي نغذيها بها، مثل صورة أو فيديو أو ملف صوتي. والطبقات الوسطى (تسمى الطبقات المخفية) تعالج البيانات، في حين تعطي طبقة المخرجات نتيجة الحوسبة (على سبيل المثال، التعرف على عناصر في صورة أو تحويل الصوت إلى نص). تُعدّ الشبكة العصبية بسيطة أو ضحلة إذا كانت تحتوي على طبقة مخفية واحدة فقط، وتكون

التعليم العميق (DEEP LEARNING (DL))

أسلوب تعلم آلة يعتمد على الشبكات العصبية التي تتكون من طبقتين داخليتين على الأقل.

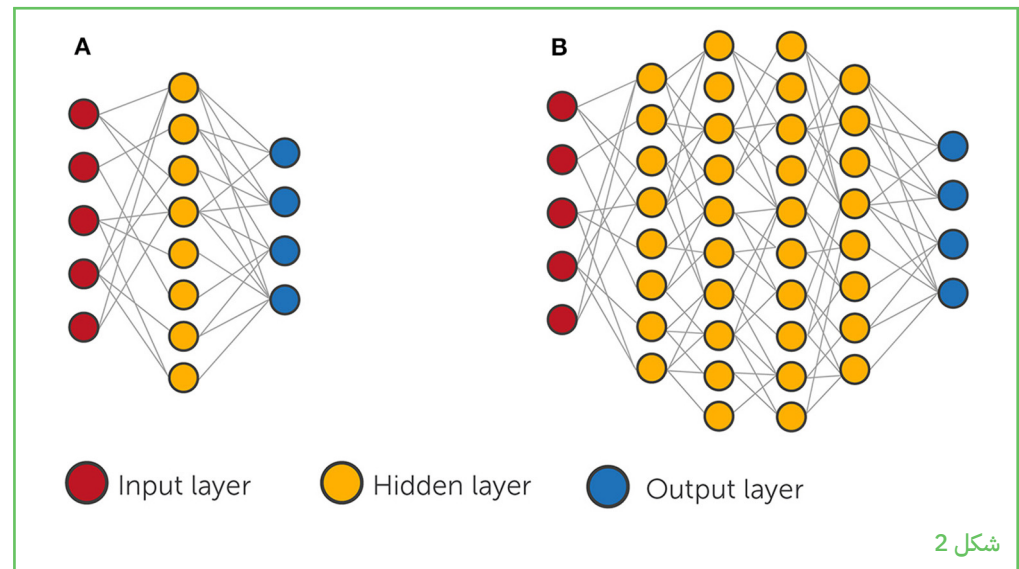
شكل 2

الشبكات العصبية الاصطناعية الشبكات العصبية الاصطناعية هي مجموعات من العصبونات الاصطناعية المنظمة في طبقات. (A) شبكة عصبية بسيطة: «طبقة مخفية» واحدة. (B) شبكة عصبية عميقة: طبقتان مخفيتان أو أكثر. يمكن للشبكات العصبية العميقة الكبيرة الاشتغال على ما يصل إلى ألف عصبون في كل طبقة وعشرات الطبقات المخفية.

شكل 3

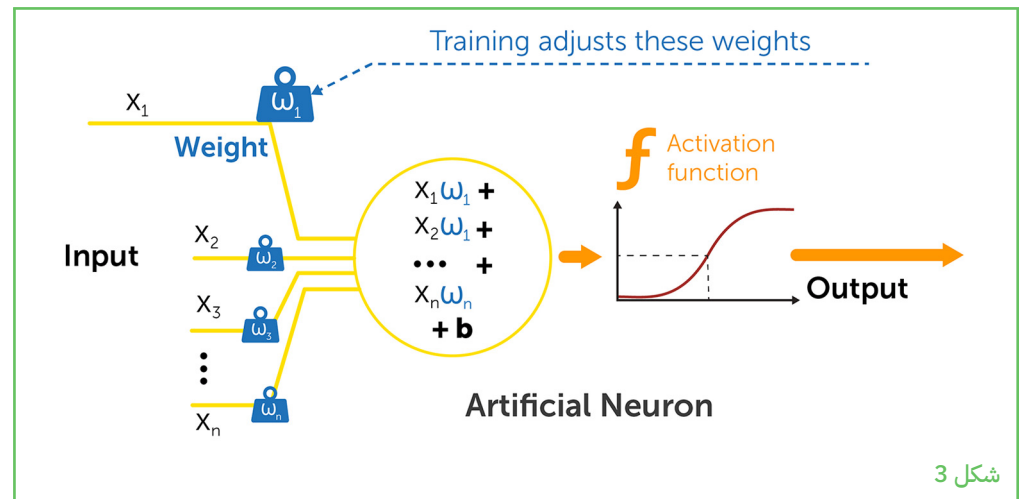
كيف تتعلم الشبكات العصبية: يستقبل كل عصبون اصطناعي في شبكة عصبية إشارات مدخلات (x_1 إلى x_n) من العصبونات الاصطناعية في الطبقة السابقة. ويضرب هذه المدخلات باستخدام مجموعة من الأوزان (w_1 إلى w_n)، ويجمعها ثم يضيف خاصية ثابتة أخرى تُسمى الانحياز (b)، والتي تحدد مدى نشاط العصبون. وفي النهاية، يُعرض هذا المجموع من خلال دالة اسمها دالة التنشيط، والتي تحدد الحاصل النهائي الذي يرسله العصبون وترسل المخرجات النهائية إلى كل العصبونات في الطبقة التالية. ويحدث التعلم في الشبكة من خلال ضبط مجموعة الأوزان والانحياز اللذين تستخدمهما العصبونات الاصطناعية في حساباتها. لمعرفة المزيد حول الحسابات في الشبكات العصبية، انظر هنا.

عميقة إذا كانت تحتوي على طبقتين مخفيتين أو أكثر. في هذه المقالة، سنركز على الشبكات العصبية العميقة التي تقوم بالتعليم العميق [5].



شكل 2

وكل عصبون اصطناعي في الشبكة العصبية العميقة يكون متصلاً بكل العصبونات الأخرى في الطبقة التالية من الشبكة، مع تفاوت قوى الاتصال، والتي تُسمى بالأوزان (الشكل 3). يستقبل كل عصبون مدخلات من كل العصبونات الأخرى في الطبقة السابقة ويقوم بعملية حسابية على أساس أوزان الاتصالات، ثم يرسل مخرجات إلى كل العصبونات الأخرى في الطبقة التالية (الشكل 3).



شكل 3

بهذه الطريقة، تتدفق المعلومات من طبقة المدخلات إلى طبقة المخرجات فيما يسمى بالانتشار الأمامي. في طبقة المخرجات، تعلن الشبكة العصبية العميقة عن النتيجة التي تراها صحيحة (على سبيل المثال، ما إذا كانت هناك قطة أم كلب في الصورة).

وهنا تأتي خطوة مهمة حيث يجب أن نتعلم شبكتنا. يعني التعلم تحديث الأوزان في الشبكة العصبية العميقة (مثل المشابك العصبية المتغيرة في الدماغ) للقيام بالمهمة بصورة أفضل في المرة التالية. فكّر في كل وزن في الشبكة العصبية العميقة مثل عقدة يمكن ضبطها لتغيير النتيجة الإجمالية، وتوجد عدة عشرات الآلاف منها. لضبط هذه العقد، تقارن الشبكة العصبية العميقة أولاً النتيجة التي حصلت عليها (تنشيط عصبونات في طبقة المخرجات، حيث يمثل كل منها نتيجة محددة) بالنتيجة الصحيحة التي كان يجب أن تحصل عليها (التنشيط المتوقع لكل عصبون في طبقة المخرجات). ثم تُحدَّث الأوزان في كل الطبقات، بدءاً من الطبقة الأخيرة إلى الطبقة المخفية الأولى لتقليل الاختلاف بين النتيجة المطلوبة والنتيجة الفعلية.

هناك طريقتان شائعتان لتدريب الشبكات العصبية العميقة. الطريقة الأولى هي التدريب الخاضع للإشراف وهو أكثر دقة، ولكنه أقل كفاءة، حيث تقوم الأنظمة بمسح العديد من عينات البيانات المصنّفة بالفعل من قبل الناس (مثل صور القطط والكلاب المصنّفة بشكل صحيح). والطريقة الثانية هي التدريب ذاتي الإشراف حيث تحاول الشبكة إعادة بناء المعلومات التي لديها بالفعل (مدخلاتها) بعد تمثيل المعلومات داخل طبقاتها، وهذا لا يتطلب أي تدخل بشري.

الخوارزمية الشائعة المُستخدمة لتحديث أوزان الشبكة العصبية العميقة تُسمى **الانتشار الخلفي** [5, 6]. ويسمح الانتشار الخلفي للشبكة العصبية العميقة بتحسين أدائها بعد كل محاولة. يستمر تحسّن الشبكات العصبية العميقة الخاضعة للتدريب حالياً بعد كل محاولة، ولكن ينخفض معدل تحسّنها، ولذلك على الرغم من توقفنا عن الاستثمار في التدريب في مرحلة معينة، فالشبكة تستمر في التحسّن. وبعد أن تتعلم الشبكة مجموعة بيانات تدريبية، تستقبل مجموعة اختبارية للتحقق من أدائها. وعندما يكون أداؤها جيداً بما فيه الكفاية للمهمة المحددة، تكون جاهزة للاستخدام في البيانات الجديدة غير المُستخدمة سابقاً.

فهم العالم بمساعدة الشبكات العصبية التلافيفية

في القشرة البصرية، وهي معالج البيانات المرئية في الدماغ، تستجيب العصبونات للامح بصرية مختلفة تُسمى المحفزات، مثل اتجاه الخطوط [7]، أو الحواف. تتكرر مجموعات من العصبونات يرصد كل منها محفزاً مختلفاً، في كل مكان في القشرة البصرية، حيث ترصد المحفزات نفسها على أجزاء مختلفة من الصورة التي نراها. يستلهم الباحثون من بنية الدماغ، حيث يطبقون البنية الرئيسية نفسها على الشبكات العصبية الاصطناعية [8]. عملت على تطوير هذه الشبكات العصبية بشكل أكبر وتم تدريبها باستخدام الانتشار الخلفي [9]، الأمر الذي قاد في النهاية إلى تطوير الشبكات العصبية التلافيفية.

كان الافتراض وراء الشبكات العصبية العميقة أن الإشارات التي نتلقاها من العالم تتكون من عناصر بسيطة تتراكم فوق بعضها لإنشاء عناصر معقدة. تُدمج المحفزات البسيطة (مثل الخطوط والحواف) في أجزاء من العناصر (مثل الأسطح المربعة

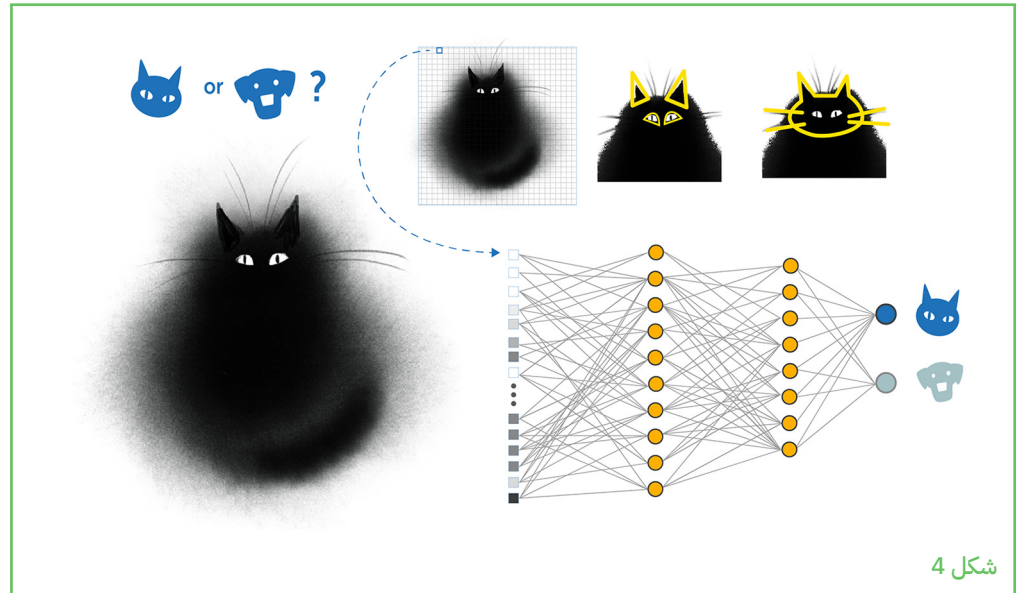
الانتشار الخلفي (BACKPROPAGATION)

خوارزمية تُعلّم تُستخدم عادةً
للتعلّم في الشبكات العصبية
العميقة.

والأرجل) ثم يُدمج كل ذلك في عناصر أكثر تعقيدًا (مثل طاولة)، بل وفئات عناصر بشكل هرمي (الشكل 4). وفي الشبكة العصبية التلافيفية، نأخذ مجموعة من العصبونات ونوصلها برقعة صغيرة في بيانات المدخلات، مثل صورة. ثم نأخذ المجموعة نفسها من العصبونات وننسخها في كل أنحاء الصورة، بحيث ننظر إلى كل جزء مختلف أو رقعة مختلفة. تحدد مخرجات العصبونات وجود أو غياب أحد الملامح في كل رقعة من الصورة. وبينما تنتشر المعلومات في كل أنحاء طبقات الشبكات العصبية التلافيفية، تصبح هذه الملامح أكثر تركيبًا أو تعقيدًا (الرؤية عرض تقديمي بديهي لآلية عمل الشبكات العصبية التلافيفية، انظر [هذا الفيديو](#)). من المدهش أننا لسنا مضطرين لإخبار الشبكات العصبية التلافيفية بالملامح التي يتوجب النظر إليها، بل ندرجها فحسب بين الأطراف باستخدام الانتشار الخلفي، ثم تكتشف بنفسها بطريقة عجيبة الملامح التي يجب استخدامها للتعرف على العنصر.

شكل 4

التعرف على الصور بمساعدة الشبكات العصبية التلافيفية: تُستخدم الشبكات العصبية التلافيفية لتطبيقات عديدة، بما في ذلك التعرف على الصور. تمثل الشبكة صورة في طبقة المدخلات كقائمة من القيم (تشير إلى لون بكسلات مختلفة في الصورة). ثم تنشر المعلومات في عناصر أكثر تعقيدًا (مثل الأذنين أو العينين أو الوجه بالكامل) إلى أن تحدد مخرجات الشبكة في النهاية العنصر الموجود في الصورة (قطة أم كلب في هذا المثال).



شكل 4

دُرِّبَت الشبكة العصبية التلافيفية التي استخدمتها على التعرف على الأرقام المكتوبة بخط اليد [9]. وقد نجحت إلى حد كبير وأُستعملت لاحقًا لقراءة شيكات في البنوك في الولايات المتحدة وأوروبا. تفيد الشبكات العصبية التلافيفية في العديد من الاستخدامات الأخرى المتعلقة بالصور (مثل التعرف البصري على الأحرف والتعرف على الوجه والمراقبة بالفيديو) بالإضافة إلى التعرف على الكلام [10, 11]. حُسِّنت الشبكات العصبية التلافيفية إلى حد كبير أداء الشبكات العصبية العميقة السابقة وأصبحت الآن جزءًا لا يتجزأ من تقنيات متطورة مثل تحليل الصور الطبية والقيادة الذاتية. بفضل مساهماتي في تطوير الشبكات العصبية التلافيفية، فزت بجائزة تورينج المرموقة في عام 2018، مشاركةً مع زميلي جيفري هينتون ويوشوا بنجيو.

هل يمكن أن تساعدنا الشبكات العصبية في فهم الدماغ؟

كما رأينا، استلهمنا تطوير الشبكات العصبية الاصطناعية من عمليات الدماغ. إذًا، هل يمكن أن تفيدنا الشبكات العصبية في فهم الدماغ نفسه؟ أرى أنها ضرورية لفهم الدماغ. فالدماغ عضو معقد للغاية وربما تكون هناك بعض المبادئ الأساسية التي تحكم قدراته، ولكنها لم تُكتشف بعد. على الرغم من أن علماء الأعصاب جمعوا كمية هائلة من البيانات التجريبية، لا توجد نظرية مؤكدة حول طريقة عمل الدماغ. لإثبات صحة نظرية حول عمل الدماغ، نحتاج إلى محاكاته في كمبيوتر وملاحظة أنه يعمل بطرق مشابهة بشكل ما لتلك الطرق في الدماغ. وإذا استطعنا بناء دماغ محوسب يعمل بشكل مشابه للدماغ الحيوي، يدل هذا على أننا حصلنا على مبادئ عمل مشتركة على الرغم من الاختلاف بين الأنظمة.

واليوم يستخدم العديد من علماء الدماغ بالفعل التعليم المعقّق والشبكات العصبية كنماذج لشرح نشاط الدماغ [12]، ولا سيما في القشرة البصرية، ولكن ذلك مهم أيضًا لتفسير كيفية معالجة الكلام والنصوص. تسمح لنا الشبكات العصبية بفحص عمليات المعلومات عن طريق تسجيل نشاط كل العصبونات واستخدام ذلك لفهم كيفية تمثيلها للبيانات. ولكن لفهم العمل الجماعي للملايين أو المليارات من هذه العناصر، يجب علينا في مرحلة ما وصف عمل الشبكة كلها على مستوى مجرد. ووجود شبكات عصبية تعمل بطرق مشابهة للدماغ سيحدث طفرة كبيرة في فهمنا الكلي للدماغ والذكاء.

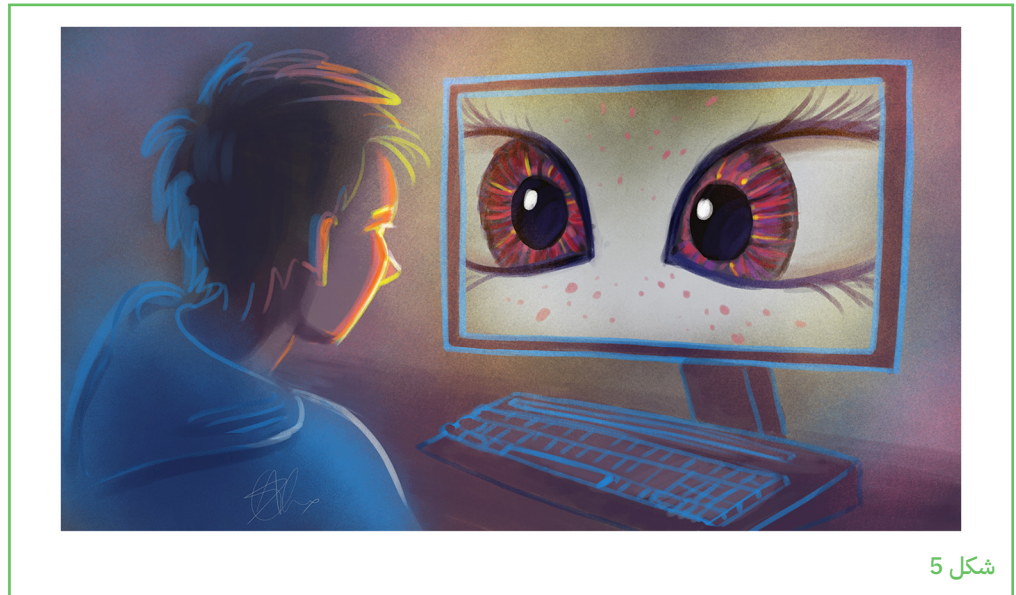
هل ستسيطر الآلات على العالم؟

أعتقد أن مشكلة تحكم الذكاء الاصطناعي قد أصبحت بمثابة "فزاعة" معاصرة، مع وجود توقعات مرعبة من أن تصبح الآلات أذكى منا وتسيطر علينا في النهاية (الشكل 5). ومع ذلك، فالذكاء لدى البشر وبالتالي الآلات لا يعني الرغبة في السيطرة على الآخر. ومن الشواغل الأخرى محاولة مواءمة سلوك الآلات الذكية مع القيم البشرية. فعلى الرغم من صعوبة "تربية" الآلات لتتصرف بالشكل السليم، يمكننا إدارتها بالطريقة نفسها التي نربي بها أطفالنا على التصرف السليم في المجتمع، وبنفس القواعد التي توجّه الأداء الاجتماعي. يمكننا تحديد الأهداف الجوهرية التي ستتبعها الآلات (تخيلها مثل "القيم الأساسية") والتي لا يجوز لها انتهاكها أو تعديلها، ما يضمن بقاء سلوكيات الآلات في تناغم مع قيمنا وأهدافنا.

وكل تكنولوجيا جديدة تجلب معها بعض العواقب غير المتوقعة، لذا يجب علينا كمجتمع أن نصحح أي آثار جانبية غير مرغوب فيها بسرعة للحد من ضررها. بعد تطوير الخدمات الإلكترونية مثل يوتيوب وفيسبوك وإنستغرام واجهنا مشكلة المحتوى غير اللائم وطورنا طرقًا للإشراف على المحتوى.

شكل 5

هل ينبغي لنا الخوف من
الآلات؟ برأيي، لا أساس للقلق
من سيطرة الآلات علينا لأنه
يمكننا ضمان كسب صداقة
الآلات، وليس عداوتها.



شكل 5

وأنا على ثقة في قدرتنا على التعامل بنجاح مع مشكلات التكنولوجيات الجديدة عند ظهورها.

إن أكثر ما يثير حماسي بشأن مستقبل الذكاء الاصطناعي هو كشف المبادئ الأساسية للذكاء. وهذا سيساعدنا على شرح المفهوم الحقيقي للذكاء، و يتيح لنا بناء الأنظمة الذكية، الأمر الذي سيوسع نطاق الذكاء البشري في النهاية. والارتقاء بفهمنا للعالم يتطلب المزيد من الذكاء، ففي مرحلة ما سنحتاج إلى أنظمة أخرى بخلاف أدمغتنا المحدودة والتي يمكننا استخدامها. للحصول على مثال مثير للاهتمام حول كيفية مساعدة الذكاء الاصطناعي لنا في فهم العالم بصورة أفضل، يمكنك قراءة [هذه المقالة](#) حول حلّ مسألة طيّ البروتين المستعصية منذ فترة طويلة.

وهناك طموح آخر خاص بالهندسة أود رؤيته يتحقق في المستقبل، وهو بناء أنظمة ذكية لمساعدتنا في حياتنا اليومية. على سبيل المثال، الروبوتات المنزلية التي ستكون بمثابة مساعدين بشريين أذكاء يديرون أشياء لا نود القيام بها ويستبعدون المعلومات غير المهمة. نسمى هذه مسألة الذكاء الاصطناعي الكامل (الأصعب) [13]، وهي تتطلب دمج العديد من القدرات والأساليب. أعمل على خوارزميات تعلم أساسية ذاتية الإشراف، أمل أن تتمكن من سدّ الفجوة بين تعلم الآلة اليوم وتعلم البشر. وأتمنى أن نتمكن من النجاح أكثر في حل مسائل الذكاء الاصطناعي الكامل وأن نعيش حياة نتمتع فيها بالمزيد من الراحة.

مواد إضافية

- صفحة Yann التعريفية.
- مشاركات Yann على Twitter.
- اطلع على [الخطاب المزيف](#) بصوت الرئيس الأمريكي السابق باراك أوباما والذي تم إنشاؤه بواسطة الذكاء الاصطناعي.

- Can Computers Understand Humor? (هل يمكن لأجهزة الكمبيوتر فهم الفكاهة؟).
- DALL-E 2 و Lexica Aperture - برامج على الإنترنت مستندة إلى الذكاء الاصطناعي مصممة لإنشاء صور من اللغة الطبيعية.

شكر وتقدير

أود شكر **أور رافايل** على إجراء المقابلة التي استند إليها هذا المقال وعلى مشاركتي في تأليفه، كما أتوجه بالشكر إلى **أليكس بيرنشتاين** على توفير الأشكال.

المراجع

1. Pfeifer, R., and Scheier, C. 2001. *Understanding Intelligence*. Cambridge: MIT Press.
2. Wolpert, D. H., and Macready, W. G. 1997. No free lunch theorems for optimization. *IEEE trans. Evolut. Computat.* 1:67–82. doi: 10.1109/4235.585893
3. Alpaydin, E. 2016. *Machine Learning: The New AI*. Cambridge: MIT press.
4. McCulloch, W. S., and Pitts, W. 1943. A logical calculus of the ideas immanent in nervous activity. *Bullet. Mathemat. Biophys.* 5:115–33.
5. LeCun, Y., Bengio, Y., and Hinton, G. 2015. Deep learning. *Nature*. 521:436–44. doi: 10.1038/nature14539
6. LeCun, Y., Touresky, D., Hinton, G., and Sejnowski, T. 1988. "A theoretical framework for back-propagation", in *Proceedings of the 1988 Connectionist Models Summer School* (Pittsburg, PA: Morgan Kaufmann), 21–28.
7. Hubel, D. H., and Wiesel, T. N. 1959. Receptive fields of single neurones in the cat's striate cortex. *J. Physiol.* 148:574. doi: 10.1113/jphysiol.1959.sp006308
8. Fukushima, K., and Miyake, S. 1982. "Neocognitron: a self-organizing neural network model for a mechanism of visual pattern recognition", in *Competition and Cooperation in Neural Nets* (Berlin: Springer), 267–285.
9. LeCun, Y., Boser, B., Denker, J., Henderson, D., Howard, R., Hubbard, W., et al. 1989. "Handwritten digit recognition with a back-propagation network", in *Advances in Neural Information Processing Systems (NIPS 1989)*, Vol. 2 (Denver, CO: Morgan Kaufmann).
10. LeCun, Y., Kavukcuoglu, K., and Farabet, C. 2010. "Convolutional networks and applications in vision", in *Proceedings of 2010 IEEE International Symposium on Circuits and Systems* (Paris: IEEE), 253–256.
11. LeCun, Y., and Bengio, Y. 1995. "Convolutional networks for images, speech, and time series", in *The Handbook of Brain Theory and Neural Networks*, ed. M. A. Arbib (MIT Press).
12. Yamins, D. L., and DiCarlo, J. J. 2016. Using goal-driven deep learning models to understand sensory cortex. *Nature Neurosci.* 19:356–365. Available online at: <https://www.nature.com/articles/nn.4244>
13. Weston, J., Bordes, A., Chopra, S., Rush, A. M., Van Merriënboer, B., Joulin, A., et al. 2015. "Towards AI-complete question answering: a set of prerequisite toy tasks", in *arXiv*.

نُشر على الإنترنت بتاريخ: 30 أبريل 2025

المحرر: Idan Segev

مرشدو العلوم: Yunchao Tang و Tomas Emmanuel Ward

الاقتباس: LeCun Y (2025) هل ستسيطر الآلات التي يمكنها التعلّم على العالم؟
Front. Young Minds. doi: 10.3389/frym.2024.1164958-ar

مُترجم ومقتبس من: LeCun Y (2024) Will Learning Machines Take Over the World? Front. Young Minds 12:1164958.
doi: 10.3389/frym.2024.1164958

إقرار تضارب المصالح: يعمل YL لدى شركة ميتا.

حقوق الطبع والنشر © 2024 © 2025 LeCun. هذا مقال مفتوح الوصول يتم توزيعه بموجب شروط ترخيص المشاركة الإبداعية **Creative Commons Attribution License (CC BY)**. يُسمح بالاستخدام أو التوزيع أو الاستنساخ في منتديات أخرى، شريطة أن يكون المؤلف (المؤلفون) الأصلي أو مالك (مالكو) حقوق النشر مقيّدًا وأن يتم الرجوع إلى المنشور الأصلي في هذه المجلة وفقًا للممارسات الأكاديمية المقبولة. لا يُسمح بأي استخدام أو توزيع أو إعادة إنتاج لا يتوافق مع هذه الشروط.

المراجعون الصغار

OISIN، العمر: 12

يحب Oisin العزف على البيانو ولعب الشطرنج ولعب ألعاب الفيديو مع أصدقائه، كما يحب الرسم وممارسة كرة القدم. وألعاب الفيديو المفضلة له هي ماين كرافت وسيطي سكاي لاينز وسيفليازيشن VI. كما يهوى قراءة كتب هاري بوتر وسلسلة ديسك ورلد التي كتبها تيري براتشيت. يعيش Oisin في أيرلندا ويدرس بالصف السادس في المدرسة الوطنية.

ZI-AN، العمر: 8

مرحبًا، اسمي Zi-An وأنا من عائلة تعمل في التدريس ورثت عنهم حب المعرفة والتعلم. ولكن مصدر سعادتي الأكبر هو أخي الصغير بالتأكيد. أحب المرح معه وجعله يضحك. تعجّبي العلوم كثيرًا، وأريد البحث في أسرار الحياة الأبدية حتى لا يكبر الأشخاص الذين أحبهم أو يموتوا أبدًا.

المؤلفون

YANN LECUN

يشغل Yann LeCun منصب نائب الرئيس ورئيس علماء الذكاء الاصطناعي في شركة ميتا، كما أنه الأستاذ الفضي في جامعة نيويورك بمعهد كورانت للعلوم الرياضية ومركز علوم البيانات. وقد حصل على دبلوم الهندسة من ESIEE (باريس) ودرجة الدكتوراة في علوم الكمبيوتر من جامعة السوربون. وعقب أبحاثه ما بعد الدكتوراة في تورونتو، انضم إلى مختبرات بيل في



عام 1988، وعُيِّن رئيسًا لمعمل معالجة الصور في أتي & تي في عام 1996. والتحق بجامعة نيويورك كأستاذ في عام 2003 وبشركة ميتا/فيسبوك في عام 2013. وتشمل اهتماماته تعلم الآلة في مجال الذكاء الاصطناعي والتصوير الحاسوبي وعلم الروبوتات وعلم الأعصاب الحوسبي. وفاز Yann بعدة جوائز مرموقة، مثل جائزة تورينج (2018)، وزمالة AAAI (2019)، وجوقة الشرف (2020)، وهو عضو في الأكاديمية الوطنية للعلوم والأكاديمية الوطنية للهندسة والأكاديمية الفرنسية للعلوم. يمارس Yann العديد من الهوايات، بما في ذلك بناء نماذج الطائرات مع عائلته وعزف الموسيقى (وعلى الأخص موسيقى الجاز حاليًا).
*yann@cs.nyu.edu

جامعة الملك عبد الله
للعلوم والتقنية
King Abdullah University of
Science and Technology



النسخة العربية مقدمة من
Arabic version provided by