

هوش در خودرو

چگونه یادگیری ماشین
دهه آینده را شکل خواهد داد

مت برجس
کمیل علی تقوی



WIRED



انتشارات راه پرداخت

برای دانلود نسخه کامل به وبسایت
فروشگاه انتشارات راه پرداخت مراجعه کنید
way2pay.shop

نسخه نمونه



The mark of
responsible forestry
FSC® C009732

سرشناسه: برجس، مت، Burgess, Matt

عنوان و نام پدیدآور: هوش مصنوعی: چگونه یادگیری ماشین، دهه آینده را شکل خواهد داد

نویسنده: مت برجس

مترجم: کمیل علی تقوی

مشخصات نشر: تهران: راه پرداخت، ۱۴۰۰

مشخصات ظاهری: ۱۲۸ ص

شابک: ۹۷۸-۶۲۲-۷۷۰۲-۰۴-۰

وضعیت فهرست نویسی: فیپا

یادداشت: عنوان اصلی: Artificial intelligence (WIRED guides): how machine...

عنوان دیگر: چگونه یادگیری ماشین، دهه آینده را شکل خواهد داد.

شناسه افزوده: علی تقوی، کمیل، ۱۳۶۳-، مترجم

شماره کتابشناسی ملی: ۷۶۶۶۱۰۵

هوش در خودروها

چگونه یادگیری ماشین
دهه آینده را شکل خواهد داد

مت برجس
کمیل علیتقوی



WIRED



عنوان: هوش مصنوعی: چگونه یادگیری ماشین، دهه آینده را شکل خواهد داد

ناشر: راه پرداخت

نویسنده: مت برجس

مترجمان: کمیل علیتقوی

ویراستار ارشد: مینا والی

ویراستار محتوایی: قاسم سرافرازی

ویراستار فنی: یلدا شایسته‌فر

بازبینی نهایی متن: رضا قربانی

صفحه‌آرا: بهناز سعیدی

ناظر چاپ: قادر شهبازی

نوبت چاپ: اول ۱۴۰۰

شمارگان: ۱۰۰۰ نسخه

شابک: ۹۷۸-۶۲۲-۷۷۰۲-۰۴-۰

تلفن: ۰۲۱-۴۴۴۳۹۶۶

دورنگار: ۸۹۷۸۴۹۰۲

ایمیل: info@way2pay.press

وبسایت: way2pay.shop

لایتوگرافی: هنر اشکان

چاپ و صحافی: واژه

همه حقوق چاپ و نشر این اثر برای «انتشارات راه پرداخت» محفوظ است. هرگونه تکثیر، انتشار و بازنویسی این اثر یا قسمتی از آن به هر شکل و شیوه (چاپی، صوتی، ویدئویی، دیجیتال و ...) بدون اجازه کتبی ناشر ممنوع است.

فروشگاه انتشارات راه پرداخت نشانی: تهران، جنت آباد جنوبی، خیابان لاله غربی، روبه‌روی پاساژ سمرقند، خیابان حدیث، کوچه حدیث دوم، پلاک ۸

۱۱	مقدمه
	بخش اول
۱۵	سال‌های نخست
	بخش دوم
۲۷	شکوفایی هوش مصنوعی
	بخش سوم
۳۷	به‌کارگیری هوش مصنوعی
	بخش چهارم
۵۵	خطرات هوش مصنوعی
	بخش پنجم
۷۵	هوش مصنوعی و دنیای کسب‌وکار
	بخش ششم
۸۷	استفاده تسلیحاتی از هوش مصنوعی
	بخش هفتم
۹۵	پاسخ‌گویی هوش مصنوعی
	بخش هشتم
۱۰۷	آینده هوش مصنوعی
۱۱۳	فهرست واژگان

[یادداشت ناشر]

وایرد نام یک ماهنامه آمریکایی است که از ژانویه ۱۹۹۳ به صورت دیجیتالی و چاپی منتشر می‌شود. موضوع این مجله چگونگی تأثیر فناوری‌های جدید و در حال رشد بر فرهنگ، اقتصاد و سیاست است. دفتر این نشریه هم در سانفرانسیسکو در ایالت کالیفرنیا قرار دارد.

این نشریه توسط خبرنگار آمریکایی لوئیس روستو و شریکش جین متکالف بنیان نهاده شد. در سال ۱۹۹۲ کوین کلی توسط لوئیس روستو استخدام شد تا به عنوان سردبیر اجرایی مجله وایرد خدمت کند.

«چرا آینده به ما نیازی ندارد» نام مقاله ایست از بیل جوی که در آوریل سال ۲۰۰۰ در مجله وایرد به چاپ رسید. او در این مقاله می‌گوید که قدرتمندترین فناوری‌های قرن بیست و یکم مانند رباتیک، مهندسی ژنتیک و فناوری نانو، انسان و گونه‌های طبیعی در معرض خطر را تهدید می‌کنند.

آینده‌ای که ۲۲ سال پیش نویسندگان وایرد از آن صحبت می‌کردند هنوز نرسیده است و هنوز هم انسان‌ها نقش مهمی بر روی زمین دارند؛ منتها یکی از آن فناوری‌هایی که در این سال‌ها این گمان درباره‌اش مطرح شده است که جایگزین انسان شود، «هوش مصنوعی» است. هرچند که نه درباره هوش توافق داریم و نه درباره هوش مصنوعی؛ ولی قبول داریم که به یکی از مهم‌ترین روندهای فناوری سال‌های اخیر

تبدیل شده است. در کلام افراد مهم کشور هم بارها بر اهمیت هوش مصنوعی تاکید شده است؛ تا امروز هم کتاب‌های مناسبی برای مدیران در زمینه شناخت هوش مصنوعی منتشر شده است؛ یکی از بهترین کتاب‌هایی که در زمینه هوش مصنوعی پیش از این منتشر شده بود هم کتاب هوش مصنوعی انتشارات هاروارد بیزینس ریویو است که آن را هم سال ۱۳۹۹ در انتشارات راه پرداخت منتشر کردیم. این کتاب متفاوت با کتاب قبلی است و روایتی جذاب از تاریخچه هوش مصنوعی و پیش‌بینی‌هایی برای آینده آن ارائه می‌دهد. این کتاب تلاش می‌کند اتفاق‌های پراکنده در دنیای هوش مصنوعی را کنار هم قرار دهد و به مخاطب دید خوبی نسبت به جایگاه این فناوری ارائه کند.

ما در گروه رسانه‌ای راه پرداخت شیفته ماهنامه وایرد هستیم و محصولات این رسانه را خودمان هم با علاقه دنبال می‌کنیم. از نظر ما وایرد نشریه مهمی است که درباره موضوعات فناوری از زاویه دید متفاوتی صحبت می‌کند و تاثیر فناوری بر حوزه‌های غیرفناوری را خیلی خوب بررسی و تحلیل می‌کند. به همین خاطر کتاب‌هایی که این رسانه منتشر کرده برای ما جذاب بود و بلافاصله بعد از انتشار این کتاب‌ها کار ترجمه و تولید را شروع کردیم که بتوانیم هم‌زمان با دنیا این کتاب‌های خوب را به مخاطب فارسی‌زبان عرضه کنیم.

تصور می‌کنیم مطالعه این کتاب دید خوبی نسبت به آینده هوش مصنوعی در اختیار مخاطبان قرار می‌دهد و افراد حرفه‌ای می‌توانند جایگاه و نسبت خود با این فناوری را دریابند. امیدواریم توانسته باشیم کار قابل قبولی را به مخاطبان خوب خود تقدیم کنیم.

رضا قربانی / انتشارات راه پرداخت

[

مقدمه

]

عصر هوش مصنوعی فرارسیده است. از ابتدای هزاره (سال ۲۰۰۰ میلادی)، فناوری نوظهوری که دهه‌ها در آزمایشگاه‌های تحقیقاتی در حال توسعه بود، وارد زندگی روزمره ما شد. شاید مانند داستان‌های علمی-تخیلی درباره هوش مصنوعی که از طریق فرهنگ عامه با آن آشنایی دارید، به نظر نرسد، اما در هر صورت، هوش مصنوعی وارد زندگی ما شده و در حال ایجاد تغییرات بزرگی در جهان است.

زمانی که تصمیم می‌گیرید این کتاب را بخوانید یا به نسخه صوتی آن گوش دهید، بی‌شک چندین بار با هوش مصنوعی تعامل خواهید داشت. شاید هر بار متوجه این تعامل نشوید؛ چراکه اکنون به قدری استفاده روزمره از هوش مصنوعی فراگیر شده که به‌ندرت متوجه آن می‌شوید.

پست‌های فیس‌بوکی که در خبرخوان (News Feed) خود آنها را لایک می‌کنید، به این دلیل به شما نشان داده می‌شوند که هوش مصنوعی این شرکت، فعالیت‌های گذشته شما را رصد کرده و می‌تواند مواردی را که احتمالاً آنها را لایک می‌کنید، پیش‌بینی کند. نتفلیکس، مصرانه فیلم‌هایی را به شما پیشنهاد می‌دهد؛ زیرا از قبل می‌داند چه چیزی را می‌پسندید. وقتی قفل تلفن همراه آیفون خود را از طریق شناسایی چهره باز می‌کنید، به این دلیل است که الگوریتم‌های شرکت اپل، ویژگی‌های چهره شما را ثبت کرده و آنها را با «اثر صورت» ذخیره شده شما، مطابقت می‌دهد.

نتایج به‌دست آمده هر جست‌وجویی که در گوگل انجام می‌دهید، مبتنی بر هوش مصنوعی است که با بررسی کل محتوای وب و دسته‌بندی هر آنچه ممکن است بیشترین ارتباط را با جست‌وجوی شما داشته باشد، به دست آمده‌اند. مسیریابی ماهواره‌ای خودرو شما، سریع‌ترین مسیرها را بر اساس درک آنی داده‌های حرکتی و دانشی که نسبت به جهان وجود دارد، پیشنهاد می‌دهد. هر بار که واژه «الکسا» را به زبان بیاورید. یا این اسپیکرهای هوشمند آمازون فکر کنند که شما کلمه بیداری (wake word) را به کار برده‌اید. هوش مصنوعی، گفته‌های شما را پردازش می‌کند و مناسب‌ترین پاسخ را بر اساس محاسباتش ارائه می‌دهد.

این فقط زندگی روزمره ما نیست که در حال دگرگون شدن است. در حوزه‌های گوناگون؛ از بهداشت و درمان گرفته تا اکتشافات فضایی، هوش مصنوعی در حال

یادگیری از انبوه داده‌هایی است که طی دهه‌ها گردآوری شده و از این داده‌ها برای کمک به کشف داروها و سیارات جدید استفاده می‌کند. در عین حال، سازمان‌ها، از شرکت‌های خصوصی گرفته تا دولت‌ها، به‌طور فزاینده‌ای در حال استفاده از هوش مصنوعی برای بهبود عملکرد خود هستند تا از بار مسئولیت اداری افرادی بکاهند که کارهای تکراری انجام می‌دهند.

ما هنوز در ابتدای راه تحول آفرین هوش مصنوعی هستیم. طی یک دهه آینده یا همین حول و حوش، هوش مصنوعی از توانایی بالایی برخوردار خواهد بود تا بسیاری از فعالیت‌های روزمره ما - مانند اتومبیل‌های خودران و تحویل روباتیک - را بهبود ببخشد و مزایای بیشتری برای جهان به ارمغان آورد. این فناوری، سردمدار مقابله با تقلب در کارت‌های اعتباری، کمک به شناسایی تصاویر بارگذاری شده درباره سوءاستفاده‌های جنسی از کودکان در اینترنت و شناسایی حملات سایبری (پیش از آنکه کسب‌وکارها را فلج کنند) خواهد بود.

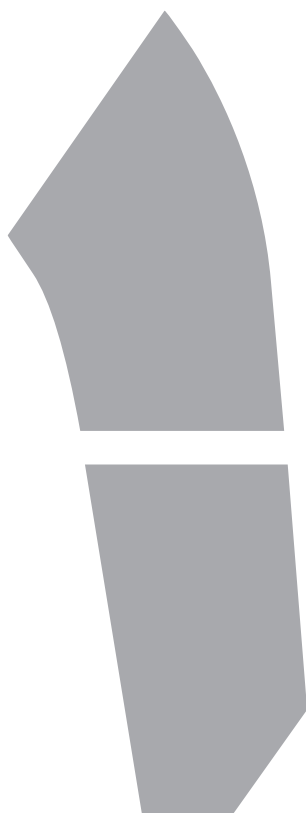
همانند سایر فناوری‌ها، دستاوردهای هوش مصنوعی نیز، تاکنون مزایا و معایبی داشته و ممکن است در آینده نیز داشته باشد. از بُعد مثبتش می‌توان گفت که هوش مصنوعی در حال مردمی و دموکراتیک شدن است. در حالی که هوش مصنوعی در گذشته فقط برای شرکت‌های دارای منابع عظیم، مناسب و مجاز بود، اکنون این اطلاعات برای هر کسی که بخواهد آنلاین باشد و ابزارهای پیاده‌سازی آن را خریداری کند، در دسترس است. در نهایت، هم شرکت‌ها و هم افراد فاقد دانش فنی، قادر خواهند بود به راحتی از هوش مصنوعی استفاده کنند. البته این نکته را هم باید در نظر داشت که هر روز بر پیچیدگی آن افزوده می‌شود.

اما مشکلات بالقوه‌ای نیز وجود دارد. در حال حاضر، اکثر تعاملات ما با سیستم‌های هوش مصنوعی نسبتاً ناچیز و بی‌اهمیت هستند؛ آنها عادی، (گاهی) بی‌وقفه و در نهایت فراموش‌شدنی هستند. یک جست‌وجوی بی‌فایده در وب یا چند بار تلاش مکرر برای باز کردن قفل تلفن، از آن دست لحظاتی هستند که بیش از چند ثانیه در خاطرم نمی‌مانند. با این حال، از این فناوری می‌توان برای جاسوسی و نظارت بر هر فعالیت استفاده کرد که این امر آزادی‌های مدنی و حقوق بشر را تهدید می‌کند و تا

حدی فقط برای همین اهداف استفاده می‌شود. بنابراین جای تعجب ندارد که نیروهای امنیتی مجذوب آن شوند یا اینکه حکومت‌ها برای تسلط بر هوش مصنوعی در حال رقابت باشند. هوش مصنوعی، فناوری‌ای است که تردیدهای زیادی را در زمینه قدرت و کنترل به وجود می‌آورد.

همچنین مواردی در زمینه نژادپرستی و تبعیض جنسیتی وجود دارد. در حال حاضر، هوش مصنوعی توسط گروه کوچکی از افراد توسعه یافته و کنترل می‌شود که اکثر آنها مردان سفیدپوستی هستند که در شرکت‌های فناوری محور فعالیت می‌کنند. چنین عدم تنوعی، برای میلیون‌ها انسانی که زندگی آنها می‌تواند تحت تأثیر سیستم‌های هوش مصنوعی قرار بگیرد، پیامدهای جدی داشته و خواهد داشت. پیاده‌سازی‌های ابتدایی حاکی از آن است که هوش مصنوعی و انواع مختلف ستم‌های اجتماعی، به همکاری با یکدیگر تمایل دارند.

در نهایت، هوش مصنوعی می‌تواند برای همیشه یک نیروی عظیم باشد و در حال حاضر، جهان را به شیوه‌های عمیقی تغییر داده است. اما دهه آینده جهت تعیین اینکه این تغییر تا کجا پیش می‌رود و در نهایت تا چه حد به نفع جامعه خواهد بود، بسیار حیاتی و سرنوشت‌ساز است. این کتاب، نه تنها چگونگی رسیدن به این مرحله را توضیح می‌دهد بلکه به بررسی این موضوع می‌پردازد که احتمالاً به کدام سو خواهیم رفت.



بخش اول
سال‌های نخست



هوش مصنوعی، آن گونه که امروزه آن را می‌شناسیم، در تعطیلات تابستانی سال ۱۹۵۶ در «کالج دارتموث» (Dartmouth) واقع در «نیوهامپشایر»، ساحل شرقی ایالات متحده آمریکا متولد شد.

طی دو ماه، گروهی از دانشمندان علوم رایانه، از جمله «جان مک‌کارتی» (John McCarthy)، «ناتانیل روچستر» (Nathaniel Rochester)، «کلود شانون» (Claude Shannon) و «ماروین مینسکی» (Marvin Minsky) که متقاعد شده بودند رایانه‌ها می‌توانند در اغلب اوقات مانند انسان‌ها رفتار کنند، گرد هم آمدند و در مورد چگونگی بهبود عملکرد ماشین‌های قدیمی بحث کردند. این چهار پژوهشگر، یک سال پیش از ملاقات‌شان در دارتموث، در مقاله‌ای نوشته بودند: «تلاش می‌کنیم بفهمیم چطور می‌توان ماشین‌ها را مجبور به استفاده از اشکال و مفاهیم انتزاعی کرد تا بتوانند مشکلاتی را حل کنند که در حال حاضر توسط انسان‌ها حل و فصل می‌شود و در نهایت خودشان را ارتقا بخشند»؛ البته حدود نیم‌دهه پیش‌تر، «آلن تورینگ» (Alan Turing)، ریاضیدان و دانشمند بریتانیایی، با ایده‌ای مبنی بر اینکه ماشین‌ها در نهایت می‌توانند خودشان به تنهایی یاد بگیرند، این موضوع را مطرح کرده بود. این دانشمندان امیدوار بودند تحقیقات آنها موجب «پیشرفتی چشمگیر» برای ماشین‌ها شود.

البته در نهایت پژوهشگران از تلاش خود ناامید شدند؛ چراکه این پروژه نتوانست همه بودجه مورد نیازشان را تأمین کند. دعوت‌شدگان در این پژوهش نتوانستند در مورد استاندارد مشترکی به توافق برسند و از آنجایی که شرکت‌کنندگان هر کدام درگیر تحقیقات خودشان بودند، پیگیری این کار با مشکل روبه‌رو شد. در هر صورت، با توجه به وضعیت نسبتاً ابتدایی رایانه‌های آن زمان، اهداف پروژه تابستانی دارتموث، غیرواقعی بود. مک‌کارتی، ۵۰ سال بعد در یادداشتی نوشت: «امیدم برای پیشرفتی چشمگیر در زمینه هوش مصنوعی در سطح انسانی، در دارتموث محقق نشد.»

اما با وجود این، پروژه تأثیرات خودش را گذاشت. پیش از انجام آن، تحقیقات مربوط به هوش مصنوعی نوظهور، اغلب در دسته‌بندی مطالعات «اتوماتا» (Automata) قرار می‌گرفت که می‌توانست به رویکردی نسبتاً ناهماهنگ منجر شود. آن چهار محقق در طرح پیشنهادی اصلی خود برای نخستین بار از عبارت «هوش مصنوعی» استفاده

کردند. مک‌کارتی نیم‌قرن بعد نوشت: «این عبارت انتخاب شد تا نظرات و عقایدمان را واضح و علنی بیان کنیم.» این استراتژی جواب داد. این عبارت جدید اکنون به صنعتی چند میلیارد دلاری تبدیل شده و اعتبار پیدایش آن نصیب خود مک‌کارتی شده است. پیشرفت اولیه کلیدی، شبکه‌های عصبی است؛ اقتباسی از یک ایده قدیمی که برای تعریف عصر مدرن هوش مصنوعی به کار می‌رود. در اوایل دهه ۱۹۴۰، دانشمندان اولین مدل‌های ریاضی سلول‌های عصبی (نورون‌ها). اجزای سازنده مغز را به وجود آوردند (که با استفاده از تکانه‌های الکتریکی (Electrical Impulses) و سیگنال‌های شیمیایی، سیگنال‌ها را در تمام مغز و سراسر سیستم عصبی بدن منتقل می‌کنند) و مدارهای الکتریکی را برای تقلید از آنها تنظیم کردند. در سال ۱۹۵۸، در پی نشست کالج دارتموث و با بودجه‌ای از سوی دفتر تحقیقات دریایی ایالات متحده آمریکا، «فرانک روزنبلات» (Frank Rosenblatt)، روان‌شناس آمریکایی، اولین شبکه عصبی مصنوعی قابل آموزش را با نام «پرسپترون» (Perceptron) ساخت که داخل ماشینی به همین نام و به اندازه یک اتاق کار می‌کرد. این شبکه شامل الگوریتمی بود که می‌توانست الگوها و اشکال ارائه‌شده به آن را تشخیص دهد.

شبکه‌های عصبی امروزی بسیار سریع‌تر و پیشرفته‌تر از پرسپترون هستند و هر کدام برای بر عهده گرفتن یک کار خاص و متفاوت طراحی می‌شوند، اما بسیاری از اصول اساسی و پایه‌ای آنها تقریباً یکسان هستند. اساساً، شبکه‌های عصبی راهی برای انجام آموزش یادگیری ماشین هستند که به موجب آن، یک رایانه با قرار گرفتن در معرض نمونه‌های پیشین یاد می‌گیرد و نهایتاً کلیه اطلاعات مورد نیاز برای تکرار آن کار را دریافت می‌کند. برای نمونه، یک شبکه عصبی را در نظر بگیرید که برای شناسایی دوچرخه در عکس‌ها در حال توسعه است. آموزش به این شبکه عصبی با نشان دادن صدها، حتی هزاران تصویر از دوچرخه انجام می‌شود که به هر یک از این تصاویر، برچسبی تحت عنوان دوچرخه زده شده تا شبکه یاد بگیرد دوچرخه چه شکلی دارد (بدیهی است وقتی دستگاهی، تصویری را «می‌بیند»، شیوه مشاهده‌اش با شیوه ما متفاوت است؛ در واقع دستگاه، عکس و متن همراه آن را به صورت یک رشته اعداد مشاهده می‌کند). اگر روند آموزش، موفقیت‌آمیز باشد، الگوریتم، قادر است دوچرخه را در عکس‌هایی که

پیش از این «ندیده» است، شناسایی کند. به طور کلی، هرچه آموزش، از طریق نمونه‌های بیشتری انجام شود (که البته زمان و هزینه بیشتری را به فرایند می‌افزاید)، نتایج آن دقیق‌تر خواهد بود. طبق همین اصل، سایر شبکه‌های عصبی به منظور ایجاد تصاویر جدید، انجام پیش‌بینی‌ها، ترجمه زبان‌ها، کنترل وسایل نقلیه خودران و کشف داروهای جدید؛ مقادیر زیادی از داده‌های عددی را تجزیه و تحلیل می‌کنند. گفتنی است شبکه‌های عصبی سیاراتی پیدا کرده‌اند، هنر و موسیقی ساخته‌اند و... و همه اینها فقط مُشتی، نمونه خروار هستند.

این سیستم‌ها تقریباً از ساختار مغز الگوبرداری شده‌اند و با اتصال شبکه‌های عصبی که حاوی میلیون‌ها یا هزاران نود پردازش داده هستند، عملکردی مشابه نورون‌ها دارند. ماشین پرسپترون، فقط از یک لایه نود استفاده می‌کرد، اما سیستم‌های مدرن می‌توانند چندین لایه. در بعضی موارد تا ۵۰ لایه. داشته باشند که داده‌ها را از طریق آنها عبور می‌دهند. در شبکه‌های عصبی با چندین لایه، هر یک از این لایه‌ها، قابلیت‌های منحصربه‌فرد خودشان را دارند و می‌توانند برای تشخیص ویژگی‌های مختلف آموزش داده شوند. چنین سیستم‌هایی، «شبکه‌های عصبی عمیق» نامیده می‌شوند و طی دو دهه اخیر کمک‌های زیادی به توسعه هوش مصنوعی در حوزه‌های گوناگون کرده‌اند.

به طور کلی، سه نوع لایه در شبکه‌های عصبی وجود دارد؛ یک لایه ورودی، لایه‌های پنهان و یک لایه خروجی. لایه ورودی جایی است که اطلاعات به شبکه عصبی منتقل می‌شود. لایه‌های پنهان، تمام پردازش‌های رایانه‌ای را انجام می‌دهند (شبکه‌های عصبی پیچیده، تعداد زیادی از این لایه‌های پنهان دارند). لایه خروجی نیز نتیجه نهایی را فراهم می‌کند. نودهای درون هر لایه به یکدیگر متصل هستند (ممکن است میلیاردها اتصال در سیستم یک شبکه عصبی بزرگ وجود داشته باشد) و به همه آنها اعدادی داده می‌شود که به عنوان «وزن» شناخته می‌شوند (وزن‌ها، کنترل می‌کنند که اتصالات بین دو نود، چقدر قوی باشند و بنابراین روی خروجی، تأثیر می‌گذارند). وزن‌ها در تمامی مراحل آموزش تغییر می‌کنند تا با خروجی‌های نهایی سازگار باشند و نتایج را باب میل سازندگان سیستم ایجاد می‌کنند.

شبکه عصبی روزنبلات می‌توانست برخی کارهای اساسی را انجام دهد، ولی از آنجایی

که فقط از یک لایه گره و یک سیستم وزنی نسبتاً ساده استفاده می‌کرد، توانایی‌های آن به شدت محدود بود. البته مردم در آن زمان چنین دیدی نسبت به پرسپترون نداشتند. در پایان جولای ۱۹۵۸، هنگامی که نیویورک تایمز گزارش داد که پرسپترون می‌تواند فقط پس از ۵۰ تلاش آموزشی، بین سمت چپ و راست تفاوت قائل شود، این ادعا مطرح شد که پرسپترون ممکن است روزی «راه برود، صحبت کند، ببیند، بنویسد، تولید مثل کند و از وجود خودش آگاه باشد». همچنین گزارش روزنبلات عنوان کرد که پرسپترون می‌تواند به عنوان بخشی از مأموریت کاوش فضایی در آینده «به سمت سیارات پرتاب شود». روزنبلات، آخرین پژوهشگر هوش مصنوعی نیست که ادعاهای جسورانه‌ای مطرح می‌کند. شایان توجه است که شبکه عصبی اولیه دیگری به نام «مادلین» (Madaline) که در سال ۱۹۵۹ توسعه داده شد، قادر به شناسایی الگوهای موجود در داده بود و کم‌اهمیت‌ترین دلیل استفاده فوری‌تر و عملی‌تر از آن نسبت به پرسپترون این بود که برای از بین بردن انعکاس صدا در خطوط تلفن - که برای تماس‌های تلفنی از راه دور استفاده می‌شد - به کار گرفته شد و آنچنان مؤثر بود که تاددها مورد استفاده قرار گرفت. زمانی که «ماروین مینسکی» (Marvin Minsky)، پیشگام هوش مصنوعی و ریاضیدانی به نام «سیمور پاپرت» (Seymour Papert)، الگوریتم پرسپترون را در کتابی با همین عنوان رد کردند محدودیت‌های فنی اولین شبکه‌های عصبی در سال ۱۹۶۹ نمایان شدند. به گفته آنها این شبکه عصبی ساده، فقط با یک لایه قادر به اجرای یکسری کارهای ابتدایی بود و خبری از توانایی موعود؛ یعنی آگاهی از وجود خودشان نبود.

کتاب مینسکی و پاپرت به دلیل کاهش اشتیاق نسبت به شبکه‌های عصبی، به طور گسترده‌ای اعتبار یافت. البته پس از انتشار آن، تحقیقات در این زمینه رفته رفته کاهش یافت، اما دهه ۱۹۶۰ به طور کلی شاهد شکوفایی تحقیقات در زمینه هوش مصنوعی و ایجاد بسیاری از ایده‌های اساسی بود که زیربنای هوش مصنوعی کنونی را تشکیل می‌دهند. دانشگاه‌های فناوری ایالات متحده آمریکا، از جمله «مؤسسه فناوری ماساچوست»، «کارنگی ملون» و «استنفورد»، تحقیق در این زمینه را آغاز کردند و شرکت‌های برجسته فناوری نیز به آنها پیوستند. بودجه این تحقیقات توسط «آژانس پروژه‌های تحقیقات پیشرفته دفاعی» (دارپا) متعلق به ارتش ایالات متحده آمریکا تأمین

شد. در همان زمان، انتشار فیلم «۲۰۰۱: ادیسه فضایی» ساخته «استنلی کوبریک» و لحن آرامش بخش شخصیت همه چیزدان و به ظاهر مفید «هال ۹۰۰۰» (Hal 9000)، آگاهی عمومی نسبت به هوش مصنوعی را افزایش داد و بذریده جهانی همزیستی ماشین‌ها با انسان‌ها را کاشت. در سال ۱۹۷۲ به نظر می‌رسید با توسعه روباتی به نام «شیکی» (Shakey) توسط مؤسسه تحقیقاتی استنفورد، یک قدم به چنین آینده‌ای نزدیک شده‌ایم. شیکی، رایانه‌ای مکعب‌شکل و دارای چرخ بود و اساساً قادر بود اشیاء را به تنهایی تشخیص داده و مسیر اطرافش را شناسایی کند و سپس توضیح دهد که چرا این کار را انجام داده است. البته ممکن بود حرکات آن (به خاطر شروع و توقف ناگهانی حرکت)، تند و نامنظم باشد (مانند نامش) و اغلب بی‌معنی بود، اما در آن زمان پیشرفته‌ترین فناوری به حساب می‌آمد.

در اوایل دهه ۱۹۸۰، پژوهشگران در حال ایجاد مجموعه‌هایی به نام «سیستم‌های خبره» بودند که می‌توانستند روند تصمیم‌گیری یک متخصص انسانی را تقلید کنند؛ البته به شرطی که وظایف مربوطه، بسیار محدود و مشخص باشد. برخی از این سیستم‌ها عملیاتی شدند. به عنوان مثال، «آمریکن اکسپرس» از یکی از همین سیستم‌های خبره برای کمک به کارکنان خود استفاده کرد تا تعیین کنند آیا مشتریان از محدودیت‌های اعتباری کارت‌هایشان فراتر رفته‌اند یا خیر. IBM این فناوری را برای کمک به مشاوران به کار گرفت تا دریابند چه میزان خدمات پس از فروش باید فروخته شود. سپس کشورهای دیگر نیز شیفته این فناوری شدند. از جمله ژاپن مبالغه‌ناگفتی در سیستم‌های هوش مصنوعی سرمایه‌گذاری کرد، اما در نهایت نتوانست نتایج چشمگیری به دست آورد. این امر نشان‌دهنده پیشرفت‌های آتی بود.

در مجموع، دهه ۱۹۷۰ و اواخر دهه ۱۹۸۰، دوران سختی برای یادگیری ماشین بود و به دوران «زمستان هوش مصنوعی» معروف شد (این نام‌گذاری از مفهوم «زمستان هسته‌ای» (Nuclear Winter) اقتباس شده است). فقدان موفقیت‌های چشمگیر در این حوزه و بهره‌گیری از سیستم‌های فنی که توان‌شان از یک حدی بالاتر نبود، باعث شد این فناوری رو به افول برود و بودجه‌های تحقیقاتی کمتری به آن اختصاص یابد. بنابراین یک احساس ناامیدی عمومی به وجود آمد؛ هوش مصنوعی نتوانسته بود وعده‌های بزرگی را

که افراد فعال در این حوزه مطرح کرده بودند، برآورده کند. در سال ۱۹۷۳، در گزارشی درباره تأثیر هوش مصنوعی که به سفارش شورای تحقیقات علوم انگلستان صورت گرفته بود، بیان شد که بیشتر پژوهشگران هوش مصنوعی و حوزه‌های مربوطه، به یک احساس مشهود ناامیدی نسبت به دستاوردهای ۲۵ سال گذشته اعتراف کرده‌اند. کسانی که در دهه‌های ۱۹۵۰ و ۱۹۶۰، وارد این حوزه شده بودند «امیدهای زیادی» به این فناوری داشتند، اما این امیدها فاصله زیادی تا محقق شدن داشتند. در این گزارش آمده است: «تاکنون در هیچ بخشی از این حوزه، اکتشافات صورت گرفته، نتوانسته‌اند تأثیر عمده‌ای را که وعده داده شده بود، عملی کنند.» «جیمز لایت‌هیل» (James Lighthill)، نویسنده گزارش، روی «ترجمه ماشینی» انگشت گذاشت تا به‌طور خاص آن را مورد انتقاد قرار دهد. او نوشت: «این مورد، «بدترین ناامیدی‌ها» را به همراه داشته است. مبالغه‌هنگفتی در این حوزه هزینه شد، اما نتایج مفید بسیار کمی به همراه داشت.»

حول و حوش همان زمانی که گزارش لایت‌هیل منتشر شد، آژانس پروژه‌های پژوهشی پیشرفته دفاعی آمریکا. دارپا. که تا آن زمان مشتاقانه از تحقیقات هوش مصنوعی پشتیبانی می‌کرد، تغییر عقیده داد و بودجه برخی تحقیقات در مورد هوش مصنوعی را قطع کرد. برخی پژوهشگران هوش مصنوعی، از این اتفاق ناراحت شدند. به عقیده آنها انتظارات غیرواقعی‌ای نسبت به توانایی‌های هوش مصنوعی وجود دارد که این انتظارات به افراد خارج از صنعت نیز سرایت کرده بود. «نیلز نیلسون» (Nils Nilsson)، پیشگام هوش مصنوعی، در کتابی که سال ۲۰۰۹ منتشر کرد، با توصیف این صنعت طی ۵۰ سال گذشته، می‌نویسد: «در اوایل دهه ۱۹۸۰، بسیاری از حامیان مالی هوش مصنوعی در دولت و صنعت، نسبت به آنچه هوش مصنوعی می‌توانست انجام دهد، انتظارات بسیار زیادی داشتند. بی‌شک می‌توان بخشی از تقصیر را بر گردن خود پژوهشگران حوزه هوش مصنوعی انداخت که انگیزه پیدا کرده بودند وعده‌های اغراق‌آمیز بدهند.» در ادامه او خاطر نشان کرد که عضویت در یک نهاد تجاری هوش مصنوعی کاهش یافت، تبلیغات در «مجله هوش مصنوعی» کم شد، برخی شرکت‌ها در حوزه هوش مصنوعی تعطیل شدند و تحقیقات در برخی از بزرگ‌ترین شرکت‌های سخت‌افزاری و نرم‌افزاری متوقف شد.

اما در ۱۱ می ۱۹۹۷، خبر قابل توجهی منتشر شد. رایانه «دیپ بلو» (Deep Blue) متعلق به IBM، «گری کاسپاروف»، استاد بزرگ و قهرمان سابق شطرنج جهان را با نتیجه ۳/۵ بر ۲/۵ شکست داد. کاسپاروف قبلاً به خود می‌بالید که هرگز در مقابل یک ماشین شکست نخواهد خورد و حتی یک سال قبل، دیپ بلو را شکست داده بود. اما این بار، از بازی آن مبهوت شد. کاسپاروف، پس از اینکه بازی آخر از سری بازی‌ها را باخت، گفت: «من یک انسان هستم. وقتی چیزی را می‌بینم که خیلی فراتر از درک من است، می‌ترسم». روی جلد شماره بعدی مجله «درون دنیای شطرنج» (Inside Chess) که پس از این اتفاق منتشر شد، نوشته شده بود: «Armageddon» (آخرالزمان).

پیروزی دیپ بلو بر کاسپاروف، نقطه عطفی را در توسعه هوش مصنوعی رقم زد و این فناوری و قابلیت‌های آن را دوباره با سرعت و شتاب بیشتری به جریان انداخت. از آن زمان به بعد هوش مصنوعی رشد روزافزونی داشته و شتاب پیشرفت‌ش طی یک دهه گذشته با کمک محاسبات ابری، سخت‌افزارهای بهبودیافته و الگوریتم‌های جدید هوشمند، بیش از پیش شده است؛ شبکه‌های عصبی به شبکه‌های عصبی عمیق تبدیل شده‌اند و یادگیری عمیق شکوفا شده است. به نحوی که امروزه تصور جهانی بدون هوش مصنوعی ممکن نیست.

هوش چیست؟

همواره نیروی محرکه هوش مصنوعی، میل به تقلید و بهبود توانایی‌های انسانی بوده است. از زمانی که این عبارت در دهه ۱۹۵۰ در دارتموث ابداع شد، تلاش‌ها معطوف به تولید دستگاهی بوده که مانند انسان‌ها فکر کند. والاترین هدف، ایجاد سیستم «هوش مصنوعی عمومی» (AGI- Artificial General Intelligence) است (که بیشتر به آن «ای جی آی» (AGI) می‌گویند).

برخی‌ها بر این باورند که هوشمندی هوش مصنوعی عمومی به اندازه انسان است و به عقیده برخی حتی این سیستم می‌تواند با هوش‌تر از انسان باشد و قادر به انجام هر کاری است که از آن خواسته می‌شود. در انجمن‌های تحقیقاتی هوش مصنوعی، بحث‌های مهمی درباره هوش مصنوعی عمومی وجود دارد؛ هم درباره شکل ظاهری آن و هم از آن

مهم‌تر، درباره اینکه آیا این سیستم واقعاً قابل دستیابی است یا خیر. در حال حاضر ما با سیستم همه‌چیزدانی که قادر به پاسخگویی به همه سؤالات باشد، بتواند تمام نیازهای بشر را تأمین کند یا حتی کنترل دنیا را بر عهده بگیرد؛ فاصله بسیار زیادی داریم، مگر اینکه پیشرفت چشمگیری در زمینه هوش مصنوعی اتفاق بیفتد. ما حتی از هوش مصنوعی که دارای هوش پایه‌ای باشد (به گونه‌ای که ما آن را درک کنیم) فاصله بسیار زیادی داریم و این موضوع به مطرح شدن یک سؤال اساسی منجر می‌شود؛ واقعاً هوش چیست؟

پاسخ کوتاه، این است که اتفاق نظر کمی در مورد این موضوع وجود دارد. وقتی نوبت به تعریف هوش می‌رسد، متخصصان حوزه هوش مصنوعی و متخصصان رشته‌های مختلف، از علوم اعصاب و روان‌شناسی گرفته تا فلسفه، اغلب به شدت اختلاف نظر دارند. وقتی نوبت به اندازه‌گیری هوش می‌رسد، هیچ‌یک از سیستم‌های فعلی، به یک اجماع جهانی نمی‌رسند (برای نمونه، آزمون‌های ضریب هوشی، از نظر بسیاری محدود و مشکل‌ساز است) و وقتی نوبت به استفاده از اصول هوش در یادگیری ماشین می‌رسد، تفسیرها متفاوت است. «هنری شولین» (Henry Shevlin)، همکار تحقیقاتی در «مرکز آینده هوش لورهلوم» (Leverhulme Centre for the Future of Intelligence) دانشگاه کمبریج می‌گوید: «فکر نمی‌کنم تعریف معناداری از هوش در هوش مصنوعی وجود داشته باشد.»

مفهوم هوش، با توجه به این واقعیت که در دوره‌های مختلف تاریخ بشر، از آن به عنوان سلاحی برای سرکوب و آسیب‌رساندن به مردم استفاده شده، پیچیده‌تر می‌شود. بین قرون هجدهم و بیستم، برخی دانشمندان ادعا کردند که نابرابری در جهان ناشی از تفاوت‌های بیولوژیکی در هوش، بین نژادهای گوناگون است؛ باهوش بودن مختص سفیدپوستان غربی قلمداد می‌شد که از آن برای پیشبرد استعمار و به عنوان مبنایی برای نژادپرستی استفاده می‌کردند. «علم نژاد»، از جنگ جهانی دوم رو به افول گذاشت، اما همان‌طور که «آنجلا سائینی» (Angela Saini) در کتابش تحت عنوان «بازگشت علم نژاد» اشاره می‌کند، هنوز هم توسط برخی گروه‌های عمدتاً راست‌گرا حمایت و تأیید می‌شود.

به گفته «کانتا دیهال» (Kanta Dihal)، پژوهشگر «مرکز آینده هوش ورهولم» در حوزه هوش مصنوعی، ایده اشتباه باهوش بودن افراد سفیدپوست همچنان پابرجاست. دیهال در حال کار روی دو پروژه مربوط به درک ما از هوش مصنوعی است؛ پروژه اول در مورد روایت‌هایی است که مردم هنگام اشاره به هوش مصنوعی استفاده می‌کنند. مورد دوم (با همکاری مدیر اجرایی مرکز «استیون کیو» (Stephen Cave)) نیز به استعمارزدایی از هوش مصنوعی می‌پردازد. دیهال می‌گوید: «کل مفهوم هوش در عبارت «هوش مصنوعی»، معمولاً امری بدیهی تلقی می‌شود، اما در واقع، معنای واژه «هوش» به شدت سیاسی و بسیار وابسته به موضوع است.» اوج هوشی که به‌طور سنتی در علوم رایانه در نظر گرفته می‌شود، ماشینی است که در بازی شطرنج متبحر باشد. اما همان‌طور که دیهال توضیح می‌دهد: «این تلقی، به اواخر قرن نوزدهم و اوایل قرن بیستم بازمی‌گردد؛ زمانی که مقیاس‌های هوش به‌منظور ایجاد و حفظ سلسله‌مراتب میان انسان‌ها در حال پایه‌ریزی بود.»

بنابراین، تعریف دقیق هوش در هوش مصنوعی، با مشکلات و هشدارهایی همراه است. با این حال، تعریف بسیار استفاده‌شده توسط محققان «شین لگ» (Shane Legg) و «مارکوس هاتر» (Marcus Hutter) را می‌توان به‌عنوان یک نقطه شروع مفید در نظر گرفت (اکنون هر دو نفر در شرکت هوش مصنوعی «دیپ مایند» متعلق به گوگل مستقر هستند که «لگ» یکی از بنیان‌گذاران آن است). آنها در سال ۲۰۰۷ این تعریف را مطرح کردند که «هوش، توانایی یک «عامل» برای دستیابی به «اهداف» یا موفقیت در طیف وسیعی از محیط‌هاست» (این دو محقق بعدها کار خود را در یک معادله ریاضی برای هوش ماشین بهبود بخشیدند). «عامل» می‌تواند یک انسان، حیوان یا ماشین باشد و «اهداف» نیز شامل حل مسئله، استدلال و پردازش سریع اطلاعات است. همگی، صفات مشترک انسانی هستند، از این رو، درک ما از هوش اساساً در خود ما استوار است.

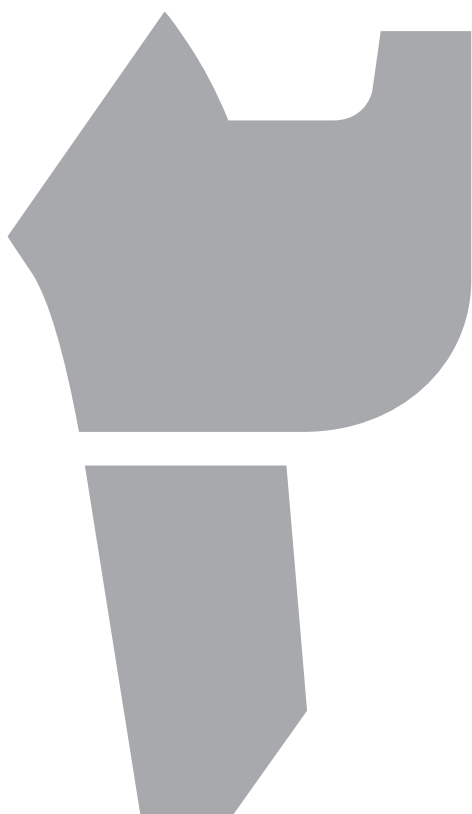
از بُعد ماشینی، اکنون چنین قابلیت‌هایی در «هوش مصنوعی محدود» (Artificial ANI - Narrow Intelligence) به جای «هوش مصنوعی عمومی» خلاصه می‌شوند. انسان‌ها قادرند از مجموعه گسترده‌ای از تجربیات بیاموزند و آنچه را آموخته‌اند، در

سناریوهای جدید به کار گیرند. هوش مصنوعی می‌تواند وظایف را در سطح انسان‌ها یا در بعضی موارد در سطحی بالاتر انجام دهد، البته در صورتی که اطلاعات کافی به آن داده شود. با این حال، خاص و محدود است. سیستم دیپ‌بلو که کاسپاروف را شکست داد یکی از بهترین بازیکنان شطرنج در تمام ادوار بود، اما به جز بازی شطرنج کار دیگری نمی‌توانست انجام دهد. حتی نمی‌توانست در بازی‌های دیگر برنده شود، چه برسد به اینکه کارهای کلی‌تر دیگری انجام دهد. حتی امروزه، هوش مصنوعی ما آن قدر پیچیده و پیشرفته نیست که بتواند بیش از یک کار انجام دهد. حتی می‌توان گفت که کودکان و حیوانات، هوش و قابلیت‌های بیشتری نسبت به بهترین سیستم‌های هوش مصنوعی ما دارند.

به عنوان مثال، کار با اشیاء را در نظر بگیرید. «چلسی فین» (Chelsea Finn)، استادیار دانشگاه استنفورد که در زمینه روباتیک فعالیت می‌کند، می‌گوید: «شما نمی‌توانید از روبات‌ها انتظار داشته باشید به همان نحوی با اشیاء سروکار داشته باشند و با آنها ارتباط برقرار کنند که یک کودک دو یا سه ساله از طریق اشیاء با جهان ارتباط برقرار می‌کند. به نظر می‌رسد که احتمالاً ما در حال از دست دادن چیزهای نسبتاً اساسی درباره هوش هستیم.» کار «فین» تا حدی روی وادار کردن روبات‌ها به تکرار آنچه انسان‌ها می‌توانند انجام دهند، متمرکز می‌شود، بنابراین آنها می‌توانند در دنیای واقعی مفید باشند. هرچه توانایی‌های یک روبات عمومی‌تر باشد، احتمال بیشتری وجود دارد که برای کارهای مختلف استفاده شود. برای تحقق یافتن این هدف، روبات‌ها باید بفهمند که چگونه اجسام در سه بُعد کار می‌کنند، چگونه می‌توان یک شیء را لمس کرد و عاقبت و پیامد جابه‌جایی یک شیء به اطراف چیست. فین می‌گوید: «فکر می‌کنم اگر ما بتوانیم به روبات‌ها اجازه دهیم مهارت‌های ارتباط با دنیای فیزیکی را داشته باشند، کارهای زیادی وجود خواهد داشت که می‌توانند برای کمک به جامعه انجام دهند. به طور مثال، برای کمک به شرایط امداد رسانی در صورت بروز فاجعه و برای کمک به مواردی که با پیری جمعیت سروکار دارند.» او اضافه می‌کند، این مشکل واقعاً دشوار است.

اگر توانایی‌های عمومی انسان، نشانه هوش باشد، پس هوش مصنوعی، راهی طولانی در پیش دارد و اگر ما از مدل هوش انسانی برای درک و به تصویر کشیدن

هوش در ماشین‌ها استفاده کنیم، این امر می‌تواند دید ما نسبت به فناوری هوش مصنوعی را تغییر دهد. در فرهنگ غرب، دیدگاه‌های ویران‌شهری درباره هوش مصنوعی بر غیرانسانی بودن آن متمرکز بوده است؛ از شخصیت‌های ۹۰۰۰ در فیلم «۲۰۰۱: یک ادیسه فضایی» گرفته تا «آوا» (Ava) در فیلم «اکس ماکینا» (Ex Machina) و همه روبات‌های فیلم «من، روبات هستم» (I Robot). دیهال خاطرنشان می‌کند که روس‌ها عموماً انتظار داشتند چت‌بات‌های ابتدایی مبتنی بر متن، مانند شخصیت آرنولد شوارتزنگر در فیلم «ترمیناتور»، از نظر احساسی دچار اختلال شوند. «توشی تاکاهاشی» (Toshie Takahashi)، استادی که در دانشگاه «واسدا» (Waseda)، در توکیو در حال مطالعه تأثیر هوش مصنوعی است، می‌گوید: «ما در ژاپن، دیدگاه مثبت‌تر و آرمان‌شهری‌تری داریم.» یک روایت محبوب در آنجا «دورائمون» (Doraemon) است؛ روایتی به شکل گربه آبی‌رنگ که چنگال و گوش ندارد و برای کمک به پسری به نام «نوبیتا نوبی» (Nobita Nobi) در زمان به عقب برمی‌گردد. شاید نتایج تحقیقات تاکاهاشی، برای ارتباط‌مان با هوش مصنوعی در آینده راهی را پیشنهاد دهد. وی نشان داده که تقریباً ۷۰ درصد جوانان در کشور، تصویری مثبت از آن دارند («مردم بیشتر در مورد تأثیر خوب اجتماعی یا مزایای اجتماعی صحبت می‌کنند، به‌ویژه برای جامعه سالخورده»)، اما ۶۰ درصد آنها نمی‌خواهند هوش مصنوعی، انسان‌نما باشد. به‌طور کلی، همان‌طور که دیهال اشاره می‌کند، رویکردی تدریجی برای ارائه هوش مصنوعی به‌عنوان چیزی که انسانی نیست و برای جامعه خطر ندارد، به وجود آمده است؛ «روایت‌هایی که هوش مصنوعی را انسان‌نما یا به‌طور کلی مجسم توصیف نمی‌کنند، افزایش یافته است.» به عبارت دیگر، ممکن است در حال حرکت به سوی آینده‌ای باشیم که هوش مصنوعی را از منظر توانایی‌های انسانی قضاوت نکنیم.



بخش دوم
[شکوفایی هوش مصنوعی]

«جف دین» (Jeff Dean) و «اندروان جی» (Andrew Ng) قصد نداشتند به هوش مصنوعی خود یاد بدهند که یک گربه چه شکلی است؛ هوش مصنوعی، خودش یاد گرفت. هنگامی که این تحقیق نوآورانه برای «یادگیری عمیق» منتشر شد، «دین» در سال ۲۰۱۲ به نیویورک تایمز گفت: «این هوش مصنوعی، اساساً خودش مفهوم گربه را ابداع کرد.»

دستاوردهای دین و ان جی، در آزمایشگاه فناوری «ایکس» گوگل نشان داد که در این ۱۵ سال، از زمانی که دیپ بلو متعلق به شرکت IBM، گری کاسپارف را شکست داد، هوش مصنوعی تا کجا پیش رفته است. این اتفاق، به برانگیختن علاقه و احیای تحقیقات در مورد هوش مصنوعی کمک کرد. چند سال بعد، هوش مصنوعی توانست در تخته‌نرد با انسان رقابت کند. سپس اولین سیستم‌های جراحی رباتیک ساخته شدند. در توانایی ماشین‌ها برای تشخیص گفتار، درک زبان نوشتاری و تشخیص ارقام دست‌نویس نیز پیشرفت‌هایی حاصل شد.

عوامل گوناگونی شاهکار دین و ان جی را ممکن کرد. با شروع هزاره جدید (سال ۲۰۰۰ میلادی)، هزینه سخت‌افزار نسبت به یک دهه پیش تر بسیار کاهش یافته بود و به نسبت سریع‌تر شده بود. به علاوه، پژوهشگران هوش مصنوعی، به این امر واقف شدند که بهترین ترانه‌های رایانه‌ای جهت استفاده، «واحدهای پردازش مرکزی» (CPU) نبودند؛ بلکه «واحدهای پردازش گرافیکی» (GPU) بودند که برای گرافیک بازی‌های رایانه‌ای به کار گرفته می‌شوند (این ترانه‌ها توسط شرکت «انویدیا» در سال ۱۹۹۹ اختراع شدند. در هر کدام از این ترانه‌ها نسبت به نسل‌های قبلی‌شان، حدود ۲۰۰ برابر پردازنده بیشتری وجود دارد که می‌توانند داده‌ها را با سرعت بیشتری پردازش کنند). در سال ۲۰۰۵ و با میلیاردها دلار کاربران اینترنت و گسترش چشمگیر تلفن‌های هوشمند، میزان داده‌های موجود برای «آموزش» شبکه‌های عصبی به‌طور تصاعدی رشد یافت.

با توجه به این پیشرفت‌ها، جای شگفتی نبود که هوش مصنوعی از دپارتمان‌های کوچک و تخصصی علوم کامپیوتر به «سیلیکون ولی» منتقل شد و در آنجا گسترش یافت. غول‌های فناوری (بیگ‌تک‌ها) مشتاقانه مشغول جذب سریع محققان هوش مصنوعی و پرداخت دستمزدهای بسیار بالا به آنها بودند تا توجه‌شان را به مجموعه داده‌های گسترده‌ای جلب کنند که متعلق به کارفرمایان جدیدشان بود. آزمایشگاه ایکس گوگل، یک نمونه برجسته

در این زمینه بود.

تیم دین وان جی به منظور آموزش به هوش مصنوعی شان که «گره چیست» یک شبکه عظیم متشکل از ۱۶ هزار پردازنده رایانه‌ای با یک میلیارد اتصال تولید کردند که روی ۱۰ میلیون تامنیل (thumbnail) ویدئویی یوتیوب به صورت تصادفی مورد استفاده قرار گرفت. پژوهشگران، از هیچ داده‌ای برای اشیای نشان داده شده (objects shown) به منظور آموزش به این سیستم استفاده نکردند. اما به تدریج، هنگامی که سیستم در معرض میلیون‌ها تصویر مشابه قرار گرفت، شروع به دسته‌بندی تک تک سلول‌های عصبی طبق ویژگی‌های اشیاء، در تامنیل‌های یوتیوب کردند. تا پایان روز سوم آموزش، سیستم یادگیری عمیق توانست گره‌ها را با دقت $74/8$ درصد، صورت انسان را با دقت $81/7$ درصد و اعضای بدن انسان را با دقت $76/7$ درصد دسته‌بندی کند. به طور کلی سیستم یادگیری عمیق موفق شد اشیاء را در ۲۰ هزار گروه مختلف دسته‌بندی کند که نسبت به سیستم‌های تشخیص تصویر قبلی، بهبود ۷۰ درصدی داشته است. این دستگاه یاد گرفته بود که اشیاء را تشخیص دهد؛ اگرچه هیچ آگاهی‌ای نسبت به اینکه آنها واقعاً چه چیزی هستند، نداشت. این، همان چیزی است که «یادگیری بدون نظارت» به خوبی انجام می‌دهد؛ دسته‌بندی اشیای مشابه در مجموعه داده‌های غیرساخت‌یافته عظیم، به گونه‌ای که به پژوهشگران کمک کند الگوهای را که به هر طریقی از دست داده‌اند، پیدا کنند.

دست‌آورد دیگری - که شاید مهم‌تر باشد - در سال ۲۰۱۲ توسط شبکه عصبی مصنوعی به دست آمد که بعداً «الکس نت» (AlexNet) نام‌گذاری شد. این شبکه توسط سه پژوهشگر در دانشگاه تورنتو به نام‌های «الکس کریژفسکی» (Alex Krizhevsky)، «ایلیا ساتسکور» (Ilya Sutskever) و «جفری هینتون» (Geoffrey Hinton). از پیشگامان دنیای هوش مصنوعی - ایجاد شد. برای نخستین بار، چنین «شبکه عصبی پیچشی» (Convolutional Neural Network) سنتی مجهز به یک GPU، توانست در مسابقه «ایمیج نت» (ImageNet) در آن سال برنده و یک شبه مشهور شود. این مسابقه در سال ۲۰۱۰ راه‌اندازی شد و دانشمندانی را درگیر می‌کرد که برای ساخت سیستمی به رقابت می‌پرداختند که قادر باشد تصاویر موجود در یک پایگاه داده رایگان، شامل میلیون‌ها تصویر دارای برچسب دستی - از خودرو و موز گرفته تا کودک - را به درستی

در هزاران دسته طبقه‌بندی کند. هر تصویر در پایگاه داده، دارای توصیفی است که توسط یک انسان افزوده شده و رقابت بر سر این بود که کدام سیستم هوش مصنوعی می‌تواند با بیشترین دقت، تصاویری را که قبلاً هرگز ندیده، شناسایی کند. عملکرد الکسنت، نسبت به رقبایش بیش از ۱۰ درصد بهتر بود و چون از پردازنده‌های گرافیکی استفاده کرد، قادر به پردازش میلیون‌ها تصویر ایمیج‌نت بود که طی پنج یا شش روز به آن نشان داده شد، در حالی که چنین پردازشی برای سیستم‌های هوش مصنوعی دارای CPU، هفته‌ها یا ماه‌ها زمان می‌برد. چنین سرعت بالایی به تیم کانادایی این امکان را داد که نتایج‌شان را مرتباً اصلاح کنند. تا سال ۲۰۱۵، پیشرفت‌های بیشتری در زمینه سرعت و دقت حاصل شده بود. برندگان ایمیج‌نت در آن سال، به دقت ۹۶ درصدی دست یافتند و برای نخستین بار نسبت به تلاش‌های انسان در شناسایی اشیاء (در تصاویر) بهتر عمل کردند. یک سال قبل‌تر، شبکه عصبی «دیپ‌فیس» (DeepFace) مربوط به تشخیص شخص در فیس‌بوک، برای بعضی از چهره‌ها به سطح دقت ۹۷/۳۵ درصد رسید؛ در حقیقت فقط ۰/۲۵ درصد کمتر از امتیاز انسان در همان آزمایش را به دست آورد. گفتنی است آموزش شبکه عصبی فیس‌بوک، با استفاده از ۴/۴ میلیون تصویر دارای برچسب (متعلق به ۴۰۳۰ کاربر این شبکه اجتماعی) انجام شد. در این زمان بود که یکی دیگر از بازیکنان اصلی در این مسابقه هوش مصنوعی، به سرتیتر اخبار تبدیل شد؛ «دیپ‌ماینند»، یک شرکت هوش مصنوعی در لندن است که در سال ۲۰۱۴ توسط گوگل با قیمت ۴۰۰ میلیون دلار خریداری شد و با اینکه استارت‌آپ شناخته‌شده‌ای نبود، اما در دنیای تحقیقات، توفان به پا کرد؛ این شرکت اعلام کرد «شبکه دیپ‌کی‌یو» (DeepQ-Network) آن یک شبکه عصبی عمیق با استفاده از یادگیری تقویتی توانسته در ۴۹ بازی کلاسیک آتاری، از جمله مهاجمان فضایی (Space Invaders) و شیوع (Breakout) پیروز شود، با وجود اینکه فقط اطلاعات محدود و مشخصی در مورد این بازی‌ها به آن داده شده است. «دمیس هاسابیس» (Demis Hassabis)، از بنیان‌گذاران این شرکت، به «وایرد» گفت: «ما در تلاش هستیم که یک مجموعه واحد از الگوریتم‌های عمومی، مانند مغز انسان را بسازیم.»

دیپ‌ماینند در مارس ۲۰۱۶، یک جهش کوانتومی روبه‌جلو داشت و با به کارگیری سیستم خود به نام «آلفاگو» توانست قهرمان جهان در بازی «گو» (Go)، با قدمت سه هزار ساله

را شکست دهد. از طرفی، این دستاورد ممکن است هم‌تراز با شکست کاسپارف توسط دیپ‌بلو در چند دهه قبل به نظر برسد، اما کاسپارف توانسته بود پایه‌پای دیپ‌بلو جلو بیاورد و در بیشتر مسابقات، به‌جز سری آخر، بازی‌ها را مساوی کند. بازی‌گو که روی یک صفحه ۱۹×۱۹ مربعی با مهره‌های سفید و سیاه بازی می‌شود، قطعاً بسیار پیچیده‌تر است (بسیاری از پژوهشگران هوش مصنوعی باورشان نمی‌شد که ماشینی بتواند بهترین بازیکنان چندین دهه جهان را شکست دهد). به‌عنوان مثال «لی سدول» (Lee Sedol)، بازیکن شماره یک بازی‌گو در جهان که ۱۸ بار برنده این بازی شده بود، بر خلاف کاسپارف، شکست سختی خورد و چهار بریک باخت.

آلفاگو در ابتدا یاد گرفت چگونه از طریق تماشای بازی انسان‌ها، بازی کند، اما با گذشت زمان، مهارت‌های تصمیم‌گیری خود را توسعه داد. در ساختار آلفاگو، چندین شبکه عصبی وجود داشت که سال‌ها پیشرفت هوش مصنوعی و قدرت مراکز داده جهانی گوگل را در خود جای داده بود. یک شبکه عصبی، حرکت بعدی در بازی را انتخاب می‌کرد، در حالی که شبکه عصبی دیگری بر اساس حرکاتی که تاکنون انجام شده بود برنده بازی را پیش‌بینی می‌کرد. در نتیجه سیستمی به وجود آمد که کارش صرفاً تقلید و تکرار هزاران حرکتی نبود که از انسان‌ها یاد گرفته؛ بلکه حرکات و سبک بازی مختص به خودش را ابداع کرد. آلفاگو در یک بازی حرکت چنان غیرمنتظره‌ای انجام دهد که هیچ انسانی تصور نمی‌کرد این حرکت را انجام دهد. «فن هیووی» (Fan Hui)، قهرمان اروپایی سه دوره بازی‌گو که این بازی را گزارش می‌کرد و با دیپ‌مایند کار کرده بود تا صدها مسابقه تمرینی در برابر آلفاگورا، پیش از پیروزی اش مقابل «لی»، انجام دهد، درباره این حرکت گفت: «این حرکت فراتر از انسان است و من هرگز انسانی را ندیده‌ام که چنین حرکتی انجام دهد. بسیار زیبا و خلاقانه بود.»

سیستم آلفاگو، بعد از مسابقه در مقابل لی، پیشرفت بیشتری کرد و از دانش بشری فراتر رفت. نسخه‌ای از این سیستم با نام «آلفاگو زیرو» (AlphaGo Zero) یاد گرفت که چگونه گوراز صفر و بدون اینکه بازی‌های قبلی انسانی را دیده باشد، بازی کند و در این فرایند آلفاگو را شکست داد. به گفته دیپ‌مایند، این نسخه توانست طی چند روز «هزاران سال دانش بشری را کنار هم جمع کند». نسخه بعدی، «آلفا زیرو» (Alpha Zero)، نه تنها قادر بود گورا بازی کند؛ بلکه شطرنج و شوگی (Shogi) را نیز بازی می‌کرد، بدون اینکه برای انجام هر بازی

متفاوت، به سازگاری نیاز داشته باشد.

در ۳۰ نوامبر ۲۰۲۰، دیپ مایند اعلام کرد که سیستم‌های هوش مصنوعی‌اش از بازی‌ها فراتر رفته‌اند تا یکی از سخت‌ترین چالش‌های زیست‌شناسی را برطرف کنند. این سیستم با نام «آلفافولد» (AlphaFold) موفق به تولید مدل‌های سه‌بعدی دقیق پروتئین‌ها شده بود؛ مواد بسیار پیچیده‌ای که در همه موجودات زنده وجود دارند. از آنجایی که پروتئین‌ها، مولکول‌های بزرگی هستند که پیش از شکل‌گیری نهایی می‌توانند سروشکل‌های بسیار متفاوتی به خود بگیرند، مدل‌سازی و پیش‌بینی آنها بسیار دشوار است. بنابراین این پیشرفت چشمگیر، بسیار مهم است. اگرچه این سیستم هنوز در روزهای ابتدایی خود قرار دارد، اما این نوید را می‌دهد که در نهایت، نسبت به فرایندهای بیولوژیکی و ژن‌هایی که سبب بیماری در انسان‌ها می‌شوند، درک و فهم بسیار بیشتری به ما ارائه خواهد داد. دمیس هاسابیس، هنگامی که نتایج آلفافولد اعلام شد، گفت: «به‌طور قطع فکر می‌کنم از نظر تأثیر در دنیای واقعی، این مهم‌ترین کاری است که انجام داده‌ایم.»

چنین پیشرفت فوق‌العاده سریعی نشان‌دهنده برخی توانایی‌های هوش مصنوعی است. در حال حاضر، این فناوری می‌تواند به کارهای منفرد به سرعت و به‌طور کارآمدی مسلط شود. اما پژوهشگران دنیای آینده را این‌گونه تصور می‌کنند که در آن یک سیستم واحد می‌تواند چندین کار را بدون نیاز به سازگاری برای هر مورد انجام دهد. اگر این چشم‌انداز چندمنظوره تحقق یابد، اساساً نحوه تعامل انسان با ماشین تغییر خواهد کرد و به مراتب بسیار بیشتر خواهد شد. همان‌یک امتیازی که لی توانست در بازی‌اش مقابل آلفاگو به دست آورد، تنها باری است که بشر تاکنون توانسته یک سیستم را شکست دهد. اما باخت کلی لی، تأثیر عمیقی روی او گذاشت. از آنجایی که حس می‌کرد دیگر هیچ چیز مانند قبل نخواهد بود، در سال ۲۰۱۹، از سری بازی‌های گو کناره‌گیری کرد و گفت: «حتی اگر من شماره یک این بازی شوم، موجودیتی وجود دارد که نمی‌توان آن را شکست داد.»

ابزار بهتر، قدرت عظیم‌تر

هوش مصنوعی با چنان سرعتی در حال توسعه و گسترش است که بسیاری از نوآوری‌هایی که تا چند سال پیش، مناسب برخی آزمایشگاه‌ها و شرکت‌های ثروتمند سیلیکون‌ولی بودند،

اکنون در دسترس اکثر ما قرار دارند. قیمت رایانه‌ها همچنان کاهش و قدرت آنها افزایش می‌یابد. اکنون می‌توانید لپ‌تاپ یا دسکتاپ پر قدرتی خریداری کنید که می‌تواند برنامه‌های هوش مصنوعی را با استفاده از سرویس‌های ابری گوگل، آمازون یا مایکروسافت اجرا کند و قادرید نحوه استفاده از آنها را به خود آموزش دهید. تمام چیزی که نیاز دارید، یک اتصال اینترنتی مناسب و چند کلیک است.

همزمان، سخت‌افزارهای این حوزه هم‌روز به‌روز پیشرفته‌تر می‌شوند. برای مثال در حالی که GPUها برای پردازش برخی برنامه‌های هوش مصنوعی ترجیح داده می‌شوند، گوگل، اینتل، «گراف‌کور» (Graphcore) و تعداد بسیار زیادی از استارت‌آپ‌ها در حال رقابت برای ساخت تراشه‌های سریع‌تر و قدرتمندتر هوش مصنوعی هستند. گوگل در سال ۲۰۱۶، از تراشه هوش مصنوعی «واحد پردازش تنسور» (TPU - Tensor Processing Unit) خود رونمایی کرد. در همین حال، «گراف‌کور» که توسط «نایجل تون» (Nigel Toon) و «سیمون نولز» (Simon Knowles) تأسیس شده، توانست برای تولید تراشه «واحد پردازش هوش» (IPU - Intelligence Processing Unit) خود بیش از ۴۰۰ میلیون دلار جذب سرمایه انجام دهد. تون، مدیرعامل شرکت گراف‌کور می‌گوید: «هنگامی که گوگل اعلام کرد در حال ساخت TPU است، لحظه بزرگی بود؛ چراکه TPU باعث شد مردم به فکر ساخت یک سخت‌افزار اختصاصی برای هوش مصنوعی بیفتند.»

IPU متعلق به شرکت گراف‌کور، از همان ابتدا نه تنها برای کار با ماشین‌های هوشمندی که امروزه وجود دارند؛ بلکه برای سرعت بخشیدن به مقیاس‌پذیری (توسعه و گسترش) بیشتر هوش مصنوعی، طراحی شده است. نسخه دوم آن - که به گفته تون، نسخه‌ای بود که همیشه قصد ساخت آن را داشت - از ۵۹/۴ میلیارد ترانزیستور (دروازه‌های الکتریکی که برای انجام محاسبات خاموش و روشن می‌شود) استفاده می‌کند که فقط هفت نانومتر از یکدیگر فاصله دارند. به گفته این شرکت، تراشه «مارک ۲» (Mark 2) می‌تواند هشت برابر عملکرد بهتری نسبت به تراشه قبلی داشته باشد. تراشه‌ها به هم متصل می‌شوند تا در دنیایی با یکدیگر ارتباط مؤثر برقرار کنند که بیش از هر زمان دیگری، برای سیستم‌های هوش مصنوعی به شدت پیچیده، داده ارسال خواهد شد. تون می‌گوید: «ما می‌دانیم این مدل‌های دانشی که قصد

داریم بسازیم، نهایتاً به مدل‌های بسیار پیچیده ختم خواهند شد. آنچه ما در تلاش هستیم انجام دهیم، ساخت قطعه‌ای سخت‌افزاری است که عملکردی پیچیده‌تر داشته باشد تا بتوانید از پراکندگی موجود در داده‌ها نهایت استفاده را ببرید. IPU اکنون وارد عملیات تجاری شده است.

بیگ‌دیتا که مایه حیات در تحول هوش مصنوعی به‌شمار می‌رود، از نظر کمیت و کیفیت نیز در حال بهبود است. «مایکل وولدریج» (Michael Wooldridge)، استاد علوم رایانه دانشگاه آکسفورد و یکی از مدیران هوش مصنوعی در مؤسسه آلن تورینگ، می‌گوید: «مردم متوجه شدند که راه پیشرفت صرفاً داشتن الگوریتم‌های هوشمندانه و قدرت محاسباتی زیاد نیست؛ بلکه آنچه شما نیاز دارید داده‌های آموزشی است که با دقت بالایی گردآوری، انتخاب و ارائه شده‌اند.» ماهواره‌ها در سراسر جهان اکنون با وضوح بسیار بالایی از سیارات عکس‌برداری می‌کنند، نسخه‌ای از تاریخ پیدایش جهان به‌صورت رایگان در ویکی‌پدیا موجود است، شهرهایمان به‌طور فزاینده‌ای مملو از حسگرهایی هستند که می‌توانند همه‌چیز را، از صدای پای افراد گرفته تا میزان پُر بودن سطل‌های زباله، ردیابی کنند. در سطح استفاده‌های شخصی نیز باید گفت هر ضربه روی تلفن همراه هوشمند شما، یک نقطه داده ایجاد می‌کند، عادات تماشای برنامه‌ها ذخیره می‌شوند، پرونده‌های پزشکی، دیجیتالی خواهند شد و افراد، در حال بارگذاری تصاویر و فیلم‌های بیشتری در وب (اینترنت) هستند.

البته دسترسی افراد سودجو به چنین داده‌هایی می‌تواند حریم خصوصی را به خطر بیندازد و امکان سوءاستفاده را فراهم کند. اما اگر در اختیار افرادی صالح قرار بگیرد، می‌توان از آن برای آموزش سیستم‌های یادگیری ماشین استفاده کرد تا به پیشرفت‌های فعالانه در جهان پیرامون ما منجر شود. محققان هم‌اکنون در حال استفاده از بیگ‌دیتا برای برنامه‌ریزی پهبادهای هستند تا بتوانند مناطق خالی از جنگل و مراتع را به‌دقت بررسی کنند و دوباره بذر بکارند. بیگ‌دیتا به هوش مصنوعی کمک کرده تا از دل صدها میلیون مولکول، آنتی‌بیوتیک‌های جدیدی کشف کند. همچنین داده‌های ناسا برای یافتن سیارات جدید پردازش می‌شوند و اقدامات مشابه دیگری نیز در جریان است.

تنها ابزارهای هوش مصنوعی نیستند که در حال افزایش هستند. خود هوش مصنوعی به‌عنوان یک صنعت، در حال بزرگ‌تر شدن و جذب سرمایه‌های بیشتری است. بر اساس

داده‌های گردآوری شده توسط شاخص هوش مصنوعی (AI Index). پروژه اجرا شده توسط دانشگاه استنفورد برای تجزیه و تحلیل وضعیت جهانی این صنعت و چگونگی تغییر آن. استارت‌آپ‌های حوزه هوش مصنوعی، ۱/۳ میلیارد دلار در سال ۲۰۱۰ جذب سرمایه داشته‌اند. در سال ۲۰۱۸، این رقم به ۴۰/۴ میلیارد دلار رسید و گرچه ممکن است همه‌گیری جهانی ویروس کرونا در سال ۲۰۲۰، به یک رکود موقت منجر شده باشد، اما جهت کلی، همچنان رو به بالاست. در سال ۲۰۱۸ بیش از سه هزار شرکت هوش مصنوعی بودجه دریافت کردند. بین سال‌های ۲۰۱۴ تا نوامبر ۲۰۱۹ بیش از ۱۵ هزار سرمایه‌گذاری با ارزشی بالغ بر ۴۰۰ هزار دلار در شرکت‌های هوش مصنوعی انجام شد که متوسط اندازه سرمایه‌گذاری، حدود ۸/۶ میلیون دلار بود. همچنین هزاران کارمند جدید توسط غول‌های فناوری استخدام شدند. بخش‌های تحقیقات هوش مصنوعی نیز افزایش یافته و بازیگران جدیدی، خارج از دایره جادویی غول‌های فناوری، جذب این حوزه شده‌اند؛ از شرکت‌های دارویی گرفته تا بانک‌ها.

هوش مصنوعی هم‌اکنون یک پروژه جهانی است. شاید هنوز بیشترین سرمایه‌گذاری‌ها تا حد زیادی توسط ایالات متحده آمریکا و چین انجام شده باشد که برای تبدیل شدن به ابرقدرت در حوزه هوش مصنوعی در حال رقابت هستند، اما سایر مناطق جهان نیز هزینه‌های کلانی خرج می‌کنند. «ولادیمیر پوتین»، رئیس‌جمهور روسیه، در سال ۲۰۱۷ گفته بود: «هر کسی که در این زمینه رهبری را در دست بگیرد، حاکم جهان خواهد شد.» میزان سرمایه‌گذاری‌های انجام شده در سال‌های ۲۰۱۸ و ۲۰۱۹ نشان داد که کشورهای مختلفی؛ از رژیم صهیونیستی گرفته تا سنگاپور و ایسلند، پیشران‌های نوآوری در هوش مصنوعی هستند. در گزارش شاخص هوش مصنوعی سال ۲۰۱۹ آمده است: «یک چیز، قطعی است؛ سیستم‌های هوش مصنوعی، به‌طور مستقیم یا غیرمستقیم، نقشی کلیدی در همه کسب‌وکارها بازی می‌کنند و اقتصاد جهانی را برای آینده‌ای قابل پیش‌بینی شکل می‌دهند.» رشد و تحول عظیمی هم در نشریات و همایش‌های مرتبط با هوش مصنوعی ایجاد شده است. گزارش‌های «شاخص هوش مصنوعی» نشان داده که طی ۲۰ سال (بین سال‌های ۱۹۹۸ تا ۲۰۱۸)، تعداد مقالات علمی هوش مصنوعی ۳۰۰ درصد رشد داشته است. امروزه چین، سالیانه به اندازه کل قاره اروپا، مجلات علمی و مقالات کنفرانسی در زمینه هوش مصنوعی منتشر می‌کند که از ایالات متحده آمریکا، به‌عنوان بزرگ‌ترین مشارکت‌کننده در حوزه هوش

مصنوعی پیشی گرفته است. همچنین بزرگ‌ترین همایش تحقیقات هوش مصنوعی به نام «سیستم‌های پردازش اطلاعات عصبی» (NeurIPS)، در سال ۲۰۱۹، ۱۳ هزار مشارکت‌کننده داشت که ۸۰۰ درصد بیشتر از سال ۲۰۱۲ بود.

امروزه، حوزه هوش مصنوعی دارای چندین زیرمجموعه است؛ از تطبیق الگو تا تشخیص صدا از طریق بینایی رایانه گرفته که به ماشین‌ها «دیدن» را می‌آموزد تا «پردازش زبان طبیعی» (NLP) که به ماشین‌ها امکان درک متن را می‌دهد. با وجود این، شاید ما هنوز در ابتدای تجاری‌سازی فناوری‌ای هستیم که در حال حاضر، در برخی مناطق بسیار فراگیر است. حتی اکنون، درک اینکه در کجا و از چه نوع هوش مصنوعی استفاده می‌شود، می‌تواند چالش‌برانگیز باشد. یک تحلیل در سال ۲۰۱۹ نتیجه‌گیری کرد که ۴۰ درصد از استارت‌آپ‌های هوش مصنوعی در اروپا، در واقع از هوش مصنوعی استفاده نمی‌کنند. اگرچه ما شاهد گسترش سریعی طی یک دهه گذشته بوده‌ایم، اما مدت زمانی طول می‌کشد تا جهان با این گسترش سریع همگام شود. مایکل وولدریج از آکسفورد استدلال می‌کند که استفاده عملی از هوش مصنوعی، هنوز در مرحله تقریباً ابتدایی قرار دارد. وی رشد آن را به رشد وب (اینترنت) تشبیه می‌کند که در سال ۱۹۸۹ اختراع شد، اما تا ۱۵ سال بعد واقعاً وارد فرایندهای کاری مردم نشده بود (و حتی امروزه بسیاری از صنایع هنوز به چاپ و امضاهای کاغذی وابسته هستند). وولدریج توضیح می‌دهد: «در دنیای واقعی، گسترش این نوع سیستم‌ها، بسیار پیچیده‌تر است.»

تا زمانی که بتوانیم از فناوری که امروزه پیرامون ماست، بهترین استفاده ممکن را داشته باشیم، راه طولانی در پیش است. اما در نهایت، هوش مصنوعی، به بخشی نامرئی از زندگی روزمره‌مان تبدیل خواهد شد. وولدریج می‌گوید: «افرادی وجود دارند و شما باید فکر کنید که چگونه آن‌ها را به طور مناسب در ادامه فرآیندهای کسب و کاریتان بگنجانید.»



بخش سوم

به کارگیری هوش مصنوعی

همه انواع هوش مصنوعی، یکسان نیستند. برخی از آنها به استانداردهای انسانی دست یافته‌اند و بعضی از آنها هنوز مسیری طولانی در پیش دارند. به عنوان مثال، سیستم‌های تشخیص گفتار مورد استفاده در اسپیکرهای هوشمند آمازون به نام «اکو» (Echo) و اسپیکرهای هوشمند خانگی گوگل، تا حدی پیشرفت کرده‌اند که برای افرادی که سؤالات ابتدایی می‌پرسند، می‌توانند مفید باشند، اما هنوز نقص‌هایی دارند. آنها اغلب دستورات ابتدایی را اشتباه می‌شنوند یا اشتباه تعبیر می‌کنند. این فناوری، در حال پیشرفت است، اما تا آن زمانی که سیستم تشخیص صدایی اختراع شود که برای مکالمات مشابه انسان‌ها مورد استفاده قرار بگیرد، فاصله داریم.

در مقابل، «بینایی رایانه‌ای» و «پردازش زبان طبیعی» دو حوزه‌ای هستند که در حال حاضر، هوش مصنوعی در آنها عالی عمل می‌کند؛ تا حدی که در برخی موارد ادعا می‌شود می‌تواند عملکرد بهتری نسبت به متخصصانی داشته باشد که تمام زندگی حرفه‌ای خودشان را وقف یک حوزه خاص کرده‌اند. هر دو، بر تلاش برای تکرار رفتار انسان متمرکز شده‌اند، یعنی «توانایی دیدن و درک جهان پیرامون آن» و «تفسیر و بازتولید صحیح گویش‌ها» و به لطف مجموعه داده‌های جدید عظیم، متشکل از میلیون‌ها و حتی تریلیون‌ها داده، هر دو، در حال برداشتن گام‌های بلند عمده‌ای روبه جلو هستند.

بینایی رایانه‌ای

«بینایی رایانه‌ای» بسیار فراتر از اتصال ساده دوربین به رایانه است. اگرچه ما کاملاً درک نمی‌کنیم چگونه مغز انسان، داده‌ها را پردازش کرده و آنچه را می‌بینیم، درک می‌کند، اما این بخش از هوش مصنوعی سعی در تکرار نتایج دارد؛ همانند چشم و مغز انسان، مراحل مختلفی درگیر هستند. یک حسگر (چشم‌ها برای یک شخص؛ نوعی دوربین برای یک ماشین) اطلاعات مربوط به آنچه را در مقابلش است به دست می‌آورد. این اطلاعات از یک مرحله تفسیر (مغز ما، یا الگوریتم‌های یک ماشین) عبور داده می‌شوند و نتیجه نهایی این است که آنچه دیده می‌شود، درک شود.

اما اینجاست که شباهت‌ها به پایان می‌رسد. انسان‌ها از همان دوران کودکی قادر به دیدن و درک اشیاء هستند؛ نیازی نیست تفاوت بین سگ و گربه به ما گفته شود. اما ماشین‌ها،

این دانش درونی را ندارند و باید داده‌های بصری را در سطح «گرانولار» (بسیار خرد) پردازش کنند. در واقع آنها یک تصویر را درست مانند یک بارکد پردازش می‌کنند. هر تصویر، بسته به وضوح و اندازه‌اش، صدها، هزاران یا حتی میلیون‌ها پیکسل خواهد داشت. یک سیستم تشخیص تصویر، همه این پیکسل‌های منفرد را به صورت اعداد می‌بیند؛ یک پیکسل سفید در یک عکس سیاه و سفید ممکن است نشان‌دهنده یک عدد پایین باشد، در حالی که به یک پیکسل سیاه می‌توان عدد بالایی اختصاص داد. این سیستم همچنین طیف‌های رنگی، شکل‌هایی که پیکسل‌ها ایجاد می‌کنند، فاصله بین این شکل‌ها و تغییرات بزرگ در اعداد را تشخیص می‌دهد. به عنوان مثال، اگر به یک سیستم تشخیص تصویر، یک «لابرادور» (نژادی از سگ‌ها) طلایی رنگ نشسته روی چمن نشان داده شود، در نقطه‌ای که تصویر از رنگ تیره چمن به پوست روشن سگ تغییر کند، سیستم تغییر بزرگی در اعداد اختصاص یافته به پیکسل‌ها نشان خواهد داد. به لطف همه پیشرفت‌های انجام شده، اکنون بینایی رایانه می‌تواند نه تنها عکس‌ها؛ بلکه فیلم‌های ویدئویی را نیز تجزیه و تحلیل کند، اشیاء را طبقه‌بندی کرده (درک کند که یک سگ، سگ است) و بشناسد (توله سگی خاص را تشخیص دهد) و اشیاء را ردیابی کند (هنگام حرکت یک سگ شکاری به اطراف، آن را دنبال کند).

اساساً، دو روش عمده به منظور آموزش ماشین برای انجام این کار وجود دارد. «یادگیری نظارت‌شده» سیستم تشخیص تصویر است که توسط تصاویری که از قبل برچسب زده می‌شوند، آموزش می‌بیند. به عبارت دیگر، برچسب‌هایی وجود دارند که روی اشیایی که سیستم سعی در تعریف آنها دارد، زده می‌شوند. زمانی که هزاران یا میلیون‌ها تصویر را تجزیه و تحلیل می‌کند، قادر است از طریق مشاهده الگوهای موجود در داده‌ها، ویژگی‌های خاص مثلاً یک سگ را بیاموزد. سپس می‌تواند یادگیری خود را برای شناسایی نژادهای مختلف سگ‌ها یا چندین شیء، از قبیل خودرو، تیر چراغ و خانه در یک تصویر گسترش دهد. با «یادگیری بدون نظارت» ماشین یاد می‌گیرد اشیایی را که ظاهری یکسان دارند گروه‌بندی کند، اما به دلیل نداشتن برچسب نمی‌داند آنها چه چیزی هستند. این نوع رویکرد سبب می‌شود الگوهای قبلاً ناشناخته داده‌ها، آشکار و قابل درک شوند و به کارهای مقدماتی کمتری در مورد داده‌ها نیاز داشته باشند، زیرا نیازی به استفاده از برچسب‌ها

نیست. از نتایج حاصل از یادگیری بدون نظارت می‌توان در آینده برای شرایط یادگیری نظارت‌شده استفاده کرد.

سیستم تشخیص تصویر، اکنون بسیار پیشرفته و در آستانه ایفای نقش مهمی در تشخیص پزشکی است. تصاویر سیاه‌وسفید گرفته‌شده حین تصویربرداری پزشکی، یکی از راه‌های تشخیصی اصلی پزشکان است. با این حال، در بسیاری از حوزه‌های پزشکی، سرعت تولید اسکن‌ها نسبت به سرعت ارزیابی‌شان، بیشتر است و از آنجایی که تشخیص به تفسیر و تخصص انسانی بستگی دارد، در این صورت ممکن است جزئیات نادیده گرفته شوند و اشتباهاتی رخ دهند. در مقابل، یادگیری عمیق می‌تواند با بررسی اسکن‌های اشعه ایکس یا ام‌آر‌آی، جزئیات ریزی را مشاهده کند که از چشم انسان‌ها دور می‌ماند. همچنین می‌تواند این کار را بسیار سریع‌تر از یک متخصص انسانی انجام دهد. در آزمایش‌ها، نشان داده شده که هوش مصنوعی قادر است سرطان (از جمله ریه، دهان، پستان، کبد، تیروئید و سرطان‌های پوستی) بیماری‌های تنفسی از جمله سل و مشکلات چشم را از طریق تصاویر موجود تشخیص دهد. برای نمونه، یک سیستم یادگیری عمیق متعلق به دانشگاه ام‌آی‌تی، با یادگیری اینکه چه الگوهایی در بافت پستان به تومورهای بدخیم منجر می‌شوند، قادر بود وجود سرطان پستان را تا پنج سال پیش از بروز کامل سرطان پیش‌بینی کند (از طریق تصاویر اسکن ۶۰ هزار بیمار یاد گرفت). در مورد دیگری، پژوهشگران بیمارستان «مورفیلدز آی» (Moorfields Eye) در لندن و شرکت دیپ‌ماینند متعلق به گوگل، در تابستان ۲۰۱۸ اعلام کردند که سیستم گوگل با استفاده از اسکن‌های «توموگرافی انسجام نوری» (OCT - optical coherence tomography)، که مدل‌های سه‌بُعدی چشم را ایجاد می‌کند، می‌تواند ۵۰ نوع مختلف از بیماری‌های چشم را به درستی (۹۴/۵ درصد از مواقع) تشخیص دهد. هوش مصنوعی (که از دو شبکه عصبی استفاده می‌کرد) این کار را با نخستین یادگیری از داده‌های تولیدشده توسط یک تیم متشکل از ۱۰ متخصص چشم‌پزشکی و بینایی‌سنجی انجام داد. این متخصصان، بیماری‌ها و علائم خاص‌شان را در اسکن تصاویر، مشخص کردند. سپس این سیستم نه‌تنها قادر بود تشخیص پزشکی را پیشنهاد دهد؛ بلکه آن را قابل اعتماد نیز می‌کرد و اسکن مربوطه را علامت‌گذاری می‌کرد

تا نشان دهد چرا چنین تصمیمی را اتخاذ کرده است.

«پیرس کین» (Pearse Keane)، یکی از جراحان بیمارستان مورفیلدز آی، در آن زمان به نشریه وایرد گفته بود: «این الگوریتم، در حد یک متخصص در زمینه تشخیص اسکن‌های OCT خوب عمل می‌کند و در تشخیص اینکه در اسکن‌های OCT، چه چیزی اشتباه است، به اندازه یک چشم‌پزشک مشاور برجسته جهانی در بیمارستان مورفیلدز آی، خوب یا حتی شاید کمی بهتر است.»

مطالعه منتشر شده در نشریه پزشکی «لنست» (Lancet) در سال ۲۰۱۹، به طور خلاصه نوید چنین فناوری‌ای را می‌داد. از میان ۳۱/۵۸۷ مقاله پژوهشی منتشر شده از سال ۲۰۱۲ در زمینه یادگیری عمیق که به تشخیص پزشکی و تصویربرداری اختصاص یافته‌اند، نویسندگان ۸۲ مورد را که اطلاعات کافی برای تجزیه و تحلیل کامل نتایج آنها وجود داشت بررسی کردند. آنها دریافتند که «عملکرد تشخیصی مدل‌های یادگیری عمیق برابر با عملکرد تشخیصی متخصصان بهداشت و درمان است». با این حال، آنها هشدار دادند: «به مطالعات بیشتری در مورد ادغام چنین الگوریتم‌هایی در دنیای واقعی نیاز است.» واقعیت، این است که تاکنون تحقیقات انجام شده بسیار امیدوارکننده بوده، اما باید اثبات شود که این نوع سیستم‌های هوش مصنوعی در آزمایش‌های بالینی - که پیش از استفاده در مقیاس انبوه، کاملاً بررسی و تکرار می‌شوند - مفید هستند یا خیر. استفاده از آنها در دنیای واقعی مستلزم آزمایش و بررسی آنها روی اسکنرهای متعدد است و باید این اطمینان حاصل شود که شرایط بیمار ایمن است و فرایندهایی را ایجاد کند که در آنها از این فناوری استفاده می‌شود. این فرایندی است که اگر به درستی و با دقت انجام شود ممکن است چندین سال زمان ببرد تا در سطح گسترده‌ای نتیجه دهد. شاید به راحتی بتوان تحت تأثیر هیجانات قرار گرفت اما نمی‌توان منکر پتانسیل موجود شد.

تشخیص پزشکی تنها حوزه‌ای نیست که در آن، تشخیص تصویر، در حال پیشرفت است. هوش مصنوعی در اپلیکیشن‌های موبایلی اکنون به کسانی که دارای مشکلات بینایی هستند می‌گوید چه اشیایی در اطراف آنها وجود دارد. در اپلیکیشن‌های «واقعیت افزوده» نیز برای قرار دادن دقیق اشیای دیجیتالی درون عکس‌ها و فیلم‌ها استفاده می‌شود. همچنین اکنون می‌تواند پیکسل‌ها را در بخش‌های بزرگ‌تری گروه‌بندی

کند. برای نمونه، اجازه می‌دهد تمام پیکسل‌هایی که یک انسان را می‌سازند، به عنوان یک شیء واحد، شناخته و به راحتی شناسایی شود.

ردیابی و شناسایی اشیاء در حال حاضر بسیار پیچیده‌تر از چند سال پیش است. یک انسان را هنگام راه رفتن در خیابان می‌توان شناسایی کرد و حضورش را می‌توان از دنیای اطرافش متمایز کرد. دوربین‌های متصل به پهپادها می‌توانند حیوانات را بشمارند یا میوه‌ها و سبزیجات در مزارع و میزان رسیده بودن آنها را ارزیابی کنند. از این دوربین‌ها سابقاً برای ارزیابی ساختمان‌ها و سازه‌های آسیب‌دیده‌ای استفاده می‌شد که برای بازرسی توسط انسان‌ها بیش از حد ناامن بودند.

شعب آمازون گو (فروشگاه‌های آزمایشی بدون صندوق‌دار متعلق به آمازون)، از ردیابی و شناسایی اشیاء، به طور گسترده‌ای استفاده می‌کنند. بینایی رایانه و حسگرهای دیگری که به کار می‌گیرند، می‌توانند مشتری را از طریق ظاهرش، شناسایی و سپس در هنگام حضور در فروشگاه او را ردیابی کنند. همچنین می‌توانند آنچه را خریداران در سبد‌هایشان می‌گذارند تشخیص دهند.

به طور کلی، بینایی رایانه می‌تواند به فروشگاه‌های خرده‌فروشی بگوید که در هر لحظه، چه تعداد مشتری در فروشگاه حضور دارند. از این فناوری در دوران همه‌گیری ویروس کرونا استفاده شد تا بررسی کند آیا مردم فاصله اجتماعی را رعایت می‌کنند. برای پخش مسابقات فوتبال نیز از آن استفاده شد تا توپ را در زمین دنبال کند و از نیاز به حضور فیزیکی نیروی انسانی گران‌قیمت در زمین فوتبال اجتناب شود و اجازه دهد مسابقات بیشتری فیلم‌برداری شود. البته این فناوری از خطا مصون نیست؛ در یک مورد، سر تاس یکی از بازیکنان در زمین باعث شد سیستم گیج شود و به جای توپ سر او را دنبال کند. در کشور استونی، یک سیستم هوش مصنوعی که داده‌های ماهواره‌ای را با اطلاعات هواشناسی ترکیب می‌کند، در حال نظارت بر مزارع کشاورزان است تا اطمینان حاصل کند که آنها تمام چمن‌زنی‌های لازم را انجام می‌دهند (که برای این کار، یارانه دولتی دریافت می‌کنند). «اوت ولزبرگ» (Ott Velsberg)، افسر ارشد داده‌های کشور استونی، می‌گوید که هنگام بازدید، بازرسان ممکن است با ادعایی از جانب کشاورزان مواجه شوند مبنی بر اینکه مزارع در ابتدای فصل چمن‌زنی شده‌اند و چمن‌ها دوباره رشد کرده‌اند یا اینکه آنها

در شرف چمن‌زنی بوده‌اند. او می‌گوید: «حتی اگر سوءظن داشتیم که کسی از بودجه دولت آن‌طور که باید استفاده نمی‌کند، باز هم هرگز هیچ پولی را پس نمی‌گرفتیم.»

اکنون، با استفاده از تصاویر جدید که هر دو روز یک بار به دست می‌آیند، نظارت پیوسته بر امور و ارسال هرگونه تذکر برای کشاورزان امکان‌پذیر است. ولزبرگ می‌گوید این پروژه در سال ۲۰۱۸، راه‌اندازی شد و در اولین سال استفاده توانست ۶۶۵ هزار یورو صرفه‌جویی در مخارج داشته باشد. او توضیح می‌دهد: «در واقع از برخی سوءاستفاده‌های احتمالی از یارانه‌های دولتی جلوگیری کرده‌ایم. شرکتی خواستار دریافت بودجه برای زمین کشاورزی‌ای شد که دقیقاً وسط یک باتلاق قرار داشت و اطراف آن، دریاچه و رودخانه بود و دسترسی به آن عملاً امکان‌پذیر نبود. تأثیرات استفاده از این فناوری فقط شامل حال نتایج به‌دست آمده از الگوریتم‌ها نمی‌شود؛ چراکه ذهنیت افراد نیز تغییر کرده است.» به لطف حضور چشمگیر و گسترده هوش مصنوعی تغییر قابل توجهی در پذیرش آن به وجود آمده است. هم‌اکنون از هوش مصنوعی ماهواره‌ای در استونی، برای هدایت کشتی‌های یخ‌شکن، شناسایی گونه‌های درختی و ثبت سوابق منابع جنگلی استفاده می‌شود.

پردازش زبان طبیعی (NLP)

یکی از حوزه‌های هوش مصنوعی که در آن پیشرفت عظیمی حاصل شده «پردازش زبان طبیعی» است. امروزه، ماشین‌ها می‌توانند با انسان‌ها ارتباط برقرار کرده و پیام‌هایی را که ما برای آن‌ها می‌فرستیم، با سرعت و دقت فوق‌بشری درک کنند و توانایی شگرفی برای بهبود خدمات مشتریان در مراکز تماس دارند، اما در کل، این فناوری پیامدهایی را نیز برای جامعه به همراه دارد. توانایی تعامل طبیعی‌تر با رایانه‌ها، و انسان‌ها، دارای پیامدهای بالقوه گسترده‌ای برای مسائل اجتماعی، فرهنگی و اقتصادی است.

در پاراگراف قبلی، به استثنای جمله اول، مابقی جملات توسط هوش مصنوعی نوشته شده، نه نویسنده کتاب. شاید جمله‌بندی آن چندان عالی و کامل نباشد، اما نشان می‌دهد که دنیای پردازش زبان طبیعی از زمان پایه‌گذاری آن در دهه ۱۹۵۰ تا به امروز چه اندازه پیشرفت کرده است. این روزها، فناوری NLP عنصر اصلی چت‌بات‌ها، سیری، الکسا و سیستم‌های تکمیل خودکار در تلفن‌هاست و همه‌چیز، از ترجمه زبان گرفته تا ایجاد

خلاصه‌ای از آثار مکتوب، را دربر می‌گیرد.

شاید هیچ حوزه‌ای در هوش مصنوعی، به اندازه NLP از بیگ‌دیتا بهره‌مند نشده باشد. اندازه مدل‌های زبانی ایجادشده توسط محققان؛ بسیار عظیم است. به عنوان مثال، پاراگراف ایجادشده توسط هوش مصنوعی در بالا، از طریق مدل زبان «مگاترون-۱۱بی» (Megatron-11b) ارائه شده که دارای ۱۱ میلیارد پارامتر است. اما این تنها قطره‌ای از اقیانوس است. در اواسط سال ۲۰۲۰، لایبراتور هوش مصنوعی مستقر در سان‌فرانسیسکو با نام «اوپن‌ای‌آی» (OpenAI)، آخرین نسخه از سیستم تولید متن خود به نام GPT-3 را ارائه کرد که شامل ۱۷۵ میلیارد پارامتر مبهوت‌کننده بود.

GPT-3، شاید پیچیده‌ترین مدل زبانی تاکنون باشد. فقط کافی است مواد مورد نیاز برای متن سیستم را آماده کنید که می‌تواند شامل چند جمله یا پاراگراف باشد. در نهایت سیستم می‌تواند به نحوی عمل کرده و متنی را آماده کند که تشخیص آن متن از متون نوشته‌شده به دست انسان، تقریباً ناممکن است. اگر درباره سبک نوشتاری‌ای که باید تولید کند، به او گفته شود، می‌تواند آن سبک را تکرار کند. GPT-3 می‌تواند همانند ویلیام شکسپیر -یا هر شخصیت ادبی دیگری- بنویسد. می‌تواند داستان‌های خلاقانه بنویسد؛ زیرا به‌طور آزمایشی برای تعدیل محتوا استفاده شده است. از GPT-3 برای کدگذاری نیز استفاده شده است. یک توسعه‌دهنده، با استفاده از این سیستم یک دستگاه طراحی صفحات وب ساخت تا هر زمان توضیحات کتبی مربوط به یک طرح خاص به آن داده شود، بتواند به‌صورت خودکار کدی برای ایجاد آن تولید کند. توسعه‌دهنده دیگری، از آن برای ایجاد یک موتور جست‌وجو استفاده کرد که می‌توانست به پرسش‌های مطرح‌شده پاسخ دهد.

«سامانیو گارگ» (Samanyou Garg)، پژوهشگری است که در حال ساخت یک برنامه پاسخ‌دهی به ایمیل با استفاده از هوش مصنوعی و سیستمی برای نوشتن نسخه بازاریابی با استفاده از GPT-3 است. او درباره تجربه‌اش در استفاده از این سیستم می‌گوید: «نتایج، نسبتاً بدیع هستند، زیرا در مورد موضوعات از قبل آگاهی دارد و در زمینه‌های بسیاری آموزش دیده است. بر اساس آنچه به‌شخصه آزمایش کرده‌ام، در نوشتن متن، واقعاً خوب است. این متون می‌تواند مقالات یا پست‌های وبلاگی یا تولید محتواها و ایده‌های

بازاریابی باشد.»

همانند تمام شبکه‌های عصبی عمیق، شیوه عملکرد سیستم تولید متن نیز با یافتن الگوها در داده‌هایی که توسط آنها آموزش دیده و ایجاد ارتباط بین کلمات و نحوه استفاده از آنها (در مورد GPT-3، از طریق یادگیری بدون نظارت) صورت می‌پذیرد. هنگامی که یک حرف را در آن تایپ می‌کنید، مانند جمله اول این بخش، در اصل حدس می‌زند که چه کلماتی باید در ادامه بیایند. البته این حدس‌ها تصادفی نیستند و بر اساس محاسبه آگاهانه محتمل‌ترین کلمات ایجاد شده‌اند.

این سیستم، خیلی کامل نیست. اگر سؤالی بی‌معنا و احمقانه از آن بپرسید، پاسخی از خودش می‌سازد. گهگاهی جملاتی را ایجاد می‌کند که معنای چندانی ندارند. البته این احتمال را نیز باید در نظر گرفت که ممکن است توسط کسانی مورد استفاده قرار بگیرد که با ایجاد اخبار جعلی و اطلاعات غلط می‌خواهند نظم جامعه را برهم بزنند. همچنین خطر بازتولید تعصبات انسانی در داده‌هایی که از آنها آموخته است، وجود دارد. «جروم پسنتی» (Jerome Pesenti)، مدیر بخش هوش مصنوعی در فیس‌بوک، نمونه‌هایی از مطالب توهین‌آمیز با محتوای نژادپرستانه و تبعیض جنسیتی را که GPT-3 منتشر کرده بود توییت کرد. این مطالب توهین‌آمیز پس از آنکه یک کلمه (مانند «یهودیان» یا هشتگ جان سیاهان اهمیت دارد) به سیستم ارائه شده بود، توسط GPT-3 منتشر شد. پسنتی می‌گوید: «ایجاد خروجی‌های نژادپرستانه و تبعیض نژادی، به‌ویژه از این طریق، نباید این چنین آسان باشد.» (البته او این‌ای، سریع‌آ راه‌حلی ارائه داد که می‌تواند از این اتفاق در آینده جلوگیری کند).

با این حال، آنچه پژوهشگران هوش مصنوعی را در مورد GPT-3 هیجان‌زده می‌کند، تغییرات بزرگ و مهم آن است. در نسخه‌های پیشین این سیستم یک‌ونیم میلیارد پارامتر و سپس ۱۱۷ میلیارد پارامتر داشت، اما در حال حاضر به ۱۷۵ میلیارد پارامتر رسیده است. در واقع هر بار که ارتقا داده می‌شد، نتایجش به‌طور گسترده‌ای بهبود می‌یافت. این مطلب، نشان می‌دهد که افزایش مقدار داده‌های ارائه‌شده به چنین سیستم‌هایی به‌سادگی می‌تواند هوش مصنوعی را به «هوش مصنوعی عمومی» (AGI) نزدیک‌تر کند.

GPT-3، تنها نمونه موجود نیست. دو مدل مهم دیگر برای پردازش زبان طبیعی، به

گوگل (به نام «برت» (BERT)) و «بایدو» (Baidu)، معادل چینی گوگل (به نام «ارنی» (ERNIE)) تعلق دارد. چنین نام‌گذاری طنزآمیزی، نشان‌دهنده رقابت و همچنین شباهت‌های این دو شرکت بزرگ فناوری است که در حال توسعه هوش مصنوعی هستند. گوگل و بایدو، هر دو، منابع و دانش عظیمی از اینترنت را در اختیار دارند و هر دو آنها با موفقیت از یادگیری عمیق برای ترجمه استفاده کرده‌اند.

تاکنون بخش عمده‌ای از تحقیقات NLP، روی زبان انگلیسی متمرکز بوده است. این بدان معنی است که مناطق وسیعی از جهان، از مدل‌ها و پیشرفت‌های NLP بهره‌ای نبرده‌اند. برای مثال، آفریقا، با وجود اندازه، اهمیت جهانی عظیمش و این واقعیت که این قاره پذیرای بیش از دو هزار زبان است، بهره‌چندانی از NLP نبرده است. خوشبختانه، پروژه «Masakhane» در تلاش است تا این امر را تغییر دهد و با کمک حدوداً ۴۰۰ محقق هوش مصنوعی از ۳۰ کشور مختلف آفریقایی، ۴۹ پروژه مختلف ترجمه از بیش از ۳۸ کشور تولید کرده و مجموعه داده‌ها و مدل‌های آموزشی جدیدی ساخته و یافته‌های تحقیقاتی خود را منتشر کرده است. در مسیر انجام این کار چالش‌هایی وجود داشت که باید برطرف می‌شد. «بوناونتور دوسو» (Bonaventure Dossou)، یکی از پژوهشگران این پروژه، می‌گوید که در ساخت یک مجموعه داده برای زبان «فون» (Fon) (یکی از زبان‌های بومی کشور «بنین» (Benin))، مجبور به ایجاد یک صفحه کلید بود تا به درستی بتوان از طریق آن تایپ کرد. او می‌گوید: «هنرمندان محلی، [اکنون] قادر به نوشتن متن ترانه‌ها به زبان فون هستند.»

شاید اشتباه باشد که فکر کنیم افزودن زبان‌های بیشتر به ابزارهای ترجمه صرفاً در خدمت منافع محلی است. در واقع، می‌تواند به نفع همه باشد. کارکنان بایدو دریافتند که با افزودن یادگیری و ترجمه زبان‌های خاص‌تر، درک هوش مصنوعی‌شان از زبان به‌طور کلی افزایش می‌یابد. سخنگوی این شرکت می‌گوید: «ترسیم هرگونه ارتباط بین زبان‌های مختلف، یک مؤلفه اساسی در این فرایند است.» از ابتدای سال ۲۰۲۰، ابزار ترجمه این شرکت، بیش از ۲۰۰ زبان را پشتیبانی و صد میلیارد حرف در روز را ترجمه می‌کند. او ادامه می‌دهد: «در حالی که یک مدل ترجمه شبکه عصبی مبتنی بر یادگیری عمیق ایجاد می‌کردیم، متوجه شدیم اگرچه زبان‌ها از نظر نمادین و دستور زبان، تفاوت‌های زیادی

دارند، اما از نظر معنایی می‌توانند بسیار شبیه به هم باشند. در بسیاری جهات؛ شناخت، درک و توصیف انسان‌ها از دنیا، علی‌رغم زبانی که صحبت می‌کنند، با یکدیگر سازگار و موافق است.»

مقدار داده‌هایی که اکنون پردازش می‌شوند، به طرز وحشتناکی زیاد هستند. برای نمونه، یک مقاله تحقیقاتی در مورد GPT-3 ارائه‌شده توسط آزمایشگاه اوپن‌ای‌آی نشان می‌دهد که تیم پروژه، با مجموعه داده «کامن کرال» (Common Crawl)، که در دسترس پژوهشگران بود، شروع کرد. این مجموعه داده، دارای حدود یک تریلیون کلمه است. اوپن‌ای‌آی، سپس آن را به منظور کیفیت فیلتر کرد و داده‌هایی از چهار مجموعه داده دیگر را در آن قرار داد که یکی از آنها، ویکی‌پدیا به زبان انگلیسی بود که ۶/۱ میلیون مقاله دارد و حاوی بیش از ۳/۶ میلیارد کلمه است.

سایر داده‌ها از کتب تاریخی - که حق نشر کپی‌رایت نداشتند - و از سراسر وب، به دست آمد. میلیون‌ها صفحه وب و متن‌هایشان کپی شدند. در واقع، هر نوع نوشتاری که در اینترنت به زبان انگلیسی با آن مواجه شده‌اید، احتمالاً در این مدل گنجانده شده است (اوپن‌ای‌آی، اطلاعات کاملی از منابعی را که برای آموزش GPT-3 استفاده کرده، فاش نکرد). اما در اینجا مشکلی وجود دارد که به همه سیستم‌های یادگیری عمیق مربوط می‌شود. پردازش چنین حجم عظیمی از داده‌ها، مستلزم صرف مقدار زیادی انرژی است. اوپن‌ای‌آی هم با گفتن اینکه مدلش، «منابع قابل توجهی» را مصرف می‌کند، وجود چنین مشکلی را تأیید کرد؛ اگرچه صحبتی درباره تأثیر بالقوه زیست‌محیطی یا هزینه آن نکرد. هزینه‌های مربوط به انرژی به‌منظور آموزش این سیستم، به‌صورت تأییدنشده تا ۱۲ میلیون دلار تخمین زده می‌شود. رقم واقعی هرچه باشد، مجموعه داده‌های بزرگ‌تری در راه هستند و بدون شک، انرژی بیشتری مصرف خواهند کرد. تحقیقات مؤسسه آلن (Allen) در حوزه هوش مصنوعی نشان داد که محاسبات مورد نیاز برای تحقیقات یادگیری عمیق بین سال‌های ۲۰۱۲ تا ۲۰۱۸، ۳۰۰ هزار برابر افزایش داشته است. «اورن اتریونی» (Oren Etzioni)، مدیرعامل این مؤسسه می‌گوید: «این مشکل، به‌طور تصاعدی در حال وخیم‌تر شدن است. مردم هنوز این مسئله را به اندازه کافی جدی نگرفته‌اند.» حال سؤالی که پیش می‌آید این است که GPT-10 چه مقدار انرژی مصرف خواهد کرد؟

در ژوئن ۲۰۱۹، محققان دانشگاه «ماساچوست آمهرست» (The University of Massachusetts Amherst)، چهار مدل زبان یادگیری عمیق (GPT-2، پرت متعلق به گوگل، ال‌مو (ELMo) متعلق به آلن ان‌ال‌پی (AllenNLP) و ترانسفورمر) را آموزش دادند. میزان مصرف انرژی این مدل‌ها با توجه به زمانی که برای آموزش نیاز داشتند، چندین برابر شد. آنها دریافتند که این مدل‌ها قادر به تولید بیش از پنج برابر میزان انتشار دی‌اکسید کربن توسط یک اتومبیل متوسط آمریکایی در تمام طول عمرش هستند. «ساشا لوجونی» (Sasha Luccioni)، پژوهشگر حوزه هوش مصنوعی و مسائل مربوط به تغییرات اقلیمی در مؤسسه «میل» (Mila) دانشگاه مونترال، می‌گوید: «آنچه در حال بهینه‌شدن است، در واقع عملکرد برخی از مجموعه داده‌هاست. محققان این دانشگاه به مدت زمان کارکرد مجموعه داده‌ها توجه نمی‌کنند و حتی بسیاری از افراد به این موضوع توجهی ندارند».

محققان هوش مصنوعی باید ذهنیت خود را در مورد تأثیرات محیطی تغییر دهند و اقداماتی در این زمینه انجام دهند. لوجونی می‌گوید: «شما قطعاً انتخاب می‌کنید که چه داده‌هایی را آموزش ببینید، مدل شما چقدر بزرگ باشد و چه مرکز داده‌ای را می‌خواهید برای آموزش انتخاب کنید.» لوجونی، به همراه همکارانش، کدی ساخته که می‌تواند به‌طور خودکار به مدل‌های هوش مصنوعی افزوده شود تا تأثیرات زیست‌محیطی هوش مصنوعی را گزارش دهند که «کدکربن» (CodeCarbon) نامیده می‌شود و میزان دی‌اکسید کربن (CO2) تولیدشده توسط منابع محاسباتی مورد نیاز در مدل‌های هوش مصنوعی را تخمین می‌زند؛ این کار را با شناسایی خودکار شبکه انرژی و سخت‌افزار مورد استفاده در آموزش انجام می‌دهد و سپس توصیه‌هایی در مورد چگونگی بهبود آن پیشنهاد می‌دهد. با بزرگ‌تر شدن داده‌های آموزش هوش مصنوعی توسعه‌دهندگان باید ردپای انرژی آن را بهبود ببخشند.

اتومبیل‌های خودران

در بهار ۲۰۰۴، ۱۵ وسیله نقلیه در صحرای موه‌اوی برای یک مسابقه منحصر به فرد گرد هم آمدند. اولین وسیله نقلیه‌ای که مسافت ۲۲۹ کیلومتری را با پیدا کردن مسیر از میان صخره‌ها، آبشارها، دره‌ها و حیات وحش غیر قابل پیش‌بینی به پایان می‌رساند، می‌توانست یک میلیون دلار برای تیم طراحی به ارمغان بیاورد. شرط انجام مسابقه چه

بود؟ هیچ یک از وسایل نقلیه مجاز به داشتن راننده نبودند. همگی باید خودشان تا خط پایان مسابقه رانندگی می کردند.

این مسابقه که توسط آژانس پروژه های تحقیقاتی پیشرفته دفاعی آمریکا (دارپا) سازمان دهی شده بود، اولین رویداد «چالش بزرگ» این گروه بود. تعدادی وسایل نقلیه به سبک فیلم «مد مکس» (Mad Max) به خط شروع می رفتند؛ خودروهای شش چرخ بزرگ، اتومبیل های شاسی بلند (اس یووی) تبدیل شده و خودروهای باگی (Buggy) همگی بدون محفظه راننده. صدها نفر نیز از نزدیک شاهد این نمایش بودند.

با شروع مسابقه، ضعف های آن معلوم شد. «برنده» موفق شد فقط ۱۱/۹ کیلومتر به گندی در مسیر حرکت کند. بعضی از وسایل نقلیه به سختی از همان خط شروع مسابقه جلوتر رفتند. بوته ها اتومبیل های خودران را گیج کردند. آنها ناگهان بی حرکت متوقف شدند، زیرا قادر به حرکت نبودند. با نرده ها برخورد کرده و روی تپه ها گیر کردند.

با این حال، این رویداد تأثیر مثبتی در تحقیقات روی خودروهای خودران داشت. یک سال بعد، پنج تیم توانستند یک مسیر مشابه ۲۱۲ کیلومتری را به طور کامل طی کنند و تیم برنده در کمتر از هفت ساعت از خط پایان گذشت. اکنون، نزدیک به دوازده بعد. میلیون ها خودرو در سراسر جهان از ویژگی های رانندگی خودکار (مانند کنترل خط و اتوپارک) استفاده می کنند؛ همچنین خودروهای کاملاً خودران، هزاران کیلومتر را در جاده های عمومی رانده اند.

«انجمن مهندسان خودرو ایالات متحده» (SAE) خودروها را بر اساس شش سطح اتوماسیون (از عدم اتوماسیون در سطح صفر تا اتوماسیون کامل در سطح پنج) رتبه بندی می کند. وسایل نقلیه سطح چهار می توانند در اکثر شرایط به طور خودران رانندگی کنند. خودروهای سطح پنج قادرند در هر شرایطی خودران کار کنند. در تمام سطوح کمتر از چهار، به راننده نیاز است و راننده باید در هر لحظه آماده و هوشیار باشد. طی مراحل توسعه، هوش مصنوعی می تواند در شبیه سازی هایی به سبک بازی های ویدئویی، آموزش ببیند، تصمیم گیری و سپس تصمیمات را دنبال کند.

فناوری و سیستم های مورد نیاز برای ایجاد خودروهای کاملاً خودران، شبیه به ساخت خودروهای نیمه خودران است که در سطوح پایین تر SAE قرار دارند. وسایل نقلیه خودران،

دارای سیستم‌های روباتیک هستند و هر سیستم روباتیک دارای سه مرحله اصلی است. «جک ویست» (Jack Weast)، معمار ارشد سیستم‌ها در بخش رانندگی خودران شرکت اینتل و معاون شرکت سیستم رانندگی خودران به نام «موبایل آی» (Mobileye) متعلق به اینتل می‌گوید: «مرحله نخست، درک کردن (یا حس کردن) است، مرحله دوم، برنامه‌ریزی و مرحله سوم، اقدام کردن است.»

بینایی رایانه نقشی بسیار کلیدی دارد. دوربین‌های مستقر در اطراف وسایل نقلیه، جهان پیرامون خود را ضبط می‌کنند و قادرند اشیاء را شناسایی کرده و تعیین کنند که آنها، چه چیزی هستند. ویست می‌گوید: «شما حسگرهایی با الگوریتم‌هایی دارید که در حال درک محیط، طبقه‌بندی اشیاء و فهم چیستی آنها و ساخت یک مدل جهانی هستند که این مدل تصویری مجازی از آنچه در جهان بیرون وجود دارد، است.»

دوربین‌ها و بینایی رایانه، تنها حسگرهای موجود نیستند. لیدار (امواج نوری) و رادار (امواج رادیویی) نیز می‌توانند مورد استفاده قرار گیرند. حسگرهای آنها، سیگنال‌هایی را ارسال می‌کنند که این سیگنال‌ها به اشیاء پیرامون خودرو، از تابلوها گرفته تا عابران پیاده، برخورد می‌کند و پیش از آنکه به خودرو برگشت داده شود، یک تصویر از محیط و اشیاء پیرامون آن ایجاد می‌کند. از آنجایی که امواج رادیویی به راحتی امواج نوری جذب نمی‌شوند می‌توانند در مسافت‌های طولانی تری حرکت کنند. نکته اصلی این است که از هر ترکیبی از حسگرها و دوربین‌ها استفاده شود به تعداد زیادی از آنها نیاز است. «پل نیومن» (Paul Newman)، بنیان‌گذار شرکت نرم‌افزاری رانندگی خودران به نام «اوکسبوتیکا» (Oxbotica)، مستقر در انگلستان می‌گوید: «شما به فراوانی و استقلال نیاز دارید. این فراوانی شامل چندین دوربین، چندین لیزر و چندین رادار می‌شود. استقلال، این است که آنها به روش‌های کاملاً متفاوتی کار می‌کنند.» فقط از طریق چنین ترکیبی می‌توان از ایمنی، اطمینان حاصل کرد. شرکت نیومن در حال آزمایش روی وسایل نقلیه خودران سطح چهار در جاده‌های انگلستان است. از وسایل نقلیه خودران این شرکت، در فرودگاه‌ها و اطراف معادن استفاده می‌شود؛ زیرا این محیط‌ها قابل پیش‌بینی هستند و به همین دلیل مسائل کمتری برای وسایل نقلیه خودران ایجاد می‌کنند. از آنجایی که در صنعت خودروسازی، اجماع عمومی وجود دارد که در آینده از وسایل نقلیه خودران در مقیاس وسیع استفاده

خواهد شد، اکنون همگی، از شرکت های فناوری گرفته تا تولیدکنندگان سنتی وسایل نقلیه، در حال توسعه چنین محصولاتی هستند (در کالیفرنیا، ۶۰ شرکت دارای مجوز برای آزمایش وسایل نقلیه خودران با حضور راننده به منظور ایمنی وجود دارد).

پس از دریافت سیگنال ها توسط خودرو، همراه با داده های دوربین های بینایی رایانه و هر ابزار نقشه برداری دیگری که در حال استفاده شدن هستند، وسیله نقلیه، آماده است که فوراً به طور مجازی تصمیم گیری کند. در حقیقت، این وسیله نقلیه، به طور مدام در حال درک کردن و برنامه ریزی حرکت بعدی خود است؛ در واقع در حال مشاهده همه اشیای پیرامون خود، از خودروها و کامیون ها گرفته تا عابران پیاده و دوچرخه سوارهاست تا بتواند از طریق ارزیابی سرعت و مسیر حرکت خود، حرکت بعدی را پیش بینی کند.

درون یک وسیله نقلیه خودران چندین الگوریتم و سیستم های هوش مصنوعی در کنار رویکردهای آماری کلاسیک غیر از هوش مصنوعی در حال فعالیت هستند. به گفته جک ویست از شرکت اینتل، در مرحله برنامه ریزی یک بلوک منطقی به نام «سیاست رانندگی» وجود دارد که بر چگونگی رفتار یک وسیله نقلیه متمرکز است. بر خلاف یادگیری نظارت شده و یادگیری بدون نظارت، سیاست رانندگی بر یادگیری تقویتی مبتنی است که شامل یک هدف خاص می شود. به الگوریتم ها گفته می شود که چیزی را به دست آورند (مثلاً مقتصدانه رانندگی کنند) و سپس سعی می کنند از طریق آزمون و خطا به آن هدف دست یابند. در صورت موفقیت، به هوش مصنوعی پاداش داده می شود. پاداش این است که به هوش مصنوعی گفته می شود رفتار صحیحی را که نشان داده بیشتر تکرار کند. ویست می گوید: «پاداش ممکن است صرفه جویی در مصرف سوخت باشد، پاداش ممکن است ایمنی باشد، پاداش ممکن است زمان رسیدن به مقصد باشد. شما در تلاش برای بهینه سازی برخی موارد هستید و این روش ابزار مناسبی برای انجام آن است.»

در راستای عملی ساختن نظریه رانندگی خودکار، شرکت «ویمو» (Waymo)، متعلق به گوگل، شرکت پیشرویی است که از سال ۲۰۰۹ در لابراتوارهای تحقیقاتی گوگل آغاز به کار کرد. طی یک دهه خودروهای خودران این شرکت بیش از ۳۰ میلیون کیلومتر در جاده های عمومی سفر کردند. اکثر کارها توسط راننده پشت فرمان (به منظور ایمنی) انجام شد؛ راننده ای که آماده بود در شرایط بحرانی کنترل خودرو را به دست بگیرد. اما از سال ۲۰۱۵،

برخی سفرها بدون حضور راننده انجام شده است. طبق گفته این شرکت، از اوایل سال ۲۰۲۰، هزار تا دو هزار رانندگی خودران در یک هفته، در منطقه‌ای به مساحت ۲۶۰ کیلومتر مربع از فونیکس آریزونا انجام شد که از این تعداد ۱۰ درصد کاملاً بدون راننده بودند. در کل، در ۹ ماه اول نخست سال ۲۰۲۰، ۱۰۵ هزار کیلومتر رانندگی کاملاً خودران انجام شد. تارسیدن به روزی که وسایل نقلیه کاملاً خودران در مقیاس تجاری مورد استفاده قرار گیرند هنوز فاصله زیادی داریم. مسائل قانونی وجود دارد که باید بر آنها غلبه کرد؛ از جمله اینکه هنگام تصادف خودروی کاملاً خودران چه کسی مقصر است؟ آیا سازنده سخت‌افزار مقصر است یا سازنده نرم‌افزار؟ مسائل عملی، وجود دارد. مثلاً وقتی ۵۰ درصد از تمامی وسایل نقلیه موجود در جاده‌ها، خودران هستند و الباقی توسط انسان‌هایی هدایت می‌شوند که از پذیرش این فناوری امتناع می‌کنند، اوضاع چگونه خواهد بود؟ همچنین مسائل اجتناب‌ناپذیری در زمینه اعتماد و ایمنی وجود دارد. آیا یک ماشین می‌تواند هر بار به درستی تصمیم بگیرد؟ واقعاً چقدر می‌تواند ایمن باشد؟

تصادفات مرگبار مربوط به خودروهای نیمه خودران شناخته شده هستند. به عنوان مثال، در مارس ۲۰۱۸، یک کارمند شرکت اپل که از یکی از خودروهای خودران شرکت تسلا استفاده می‌کرد پس از آنکه خودروی او به یک مانع سیمانی برخورد کرد، جانش را از دست داد. «هیئت ملی ایمنی حمل و نقل ایالات متحده» با انجام تحقیقات سرانجام راننده وسیله نقلیه را مقصر دانست. نتیجه‌گیری کرد که در هنگام تصادف، راننده در حال بازی کردن بوده و هنگامی که سیستم به کمک نیاز داشته، دست‌هایش روی فرمان خودرو نبوده است. همچنین تصمیم‌گیری کرد که «سیستم برخورد با مانع» شرکت تسلا مانع مسبب تصادف را که خودرو با آن برخورد کرده، شناسایی نکرده و اینکه به‌طور مؤثری بر درگیر کردن راننده با جاده نظارت نداشته است.

ویمو متعلق به گوگل، جزئیاتی را درباره تعداد تصادفات مربوط به رانندگی وسایل نقلیه خودران خود در طول سفرهایی بالغ بر ۹/۸ میلیون کیلومتر منتشر کرد (۱۰۵ هزار کیلومتر از آن، همان‌طور که قبلاً ذکر شد، بدون حضور راننده بود). این یک مجموعه داده نسبتاً کوچک است، اما با وجود این، بینش جالبی در مورد نوع چالش‌هایی که وسایل نقلیه خودران با آن روبه‌رو هستند، فراهم می‌کند. به‌طور کلی، به نظر می‌رسد ۴۷ برخورد شدید

یا تماس رخ داده که ۱۸ مورد آن واقعاً اتفاق افتاده است. بر اساس شبیه‌سازی‌های بعدی ویمو، در صورتی که پای دخالت انسانی وسط نبود، ممکن بود آن ۲۹ حادثه دیگر نیز اتفاق بیفتند. در ۱۶ مورد از این تصادفات خودروهای ویمو از عقب آسیب دیده بودند و زمانی این اتفاق افتاده بود که این خودروها پشت چراغ راهنمایی و رانندگی توقف کرده بودند. در حوادث دیگر، افراد دوچرخه‌سوار یا پیاده به سمت این خودروهای خودران رفته بودند. در موارد دیگر، افراد به علت نقض قوانین رانندگی با خودروهای خودران تصادف می‌کردند. برای مثال، یک ماشین غیر خودران به دلیل اینکه در مسیر مخالف رانندگی می‌کرد با یک خودروی خودران تصادف سنگینی کرد. تنها در یک مورد از این تصادفات، خودروی ویمو، مقصر شناخته شد. تیم‌های ویمو در یک مقاله تحقیقاتی درباره تصادفات آن نوشتند: «تقریباً همه حوادث واقعی و شبیه‌سازی شده شامل یک یا چند مورد نقض قوانین رانندگی یا ناشی از بی‌احتیاطی توسط عامل دیگری می‌شد. وسایل نقلیه خودران، در آینده، جاده‌ها را با رانندگان انسانی شریک خواهند شد و تعداد قابل توجهی از تصادفات ناشی از خطاهای راننده انسانی که اجتناب‌ناپذیر نیز هستند، در این دوره پیش‌بینی می‌شوند.» به نظر می‌رسد چنین مسائلی، نظر طرفداران فناوری رانندگی خودران را مبنی بر ایمن‌تر بودن این خودروها نسبت به حمل‌ونقل عمومی، تأیید می‌کند. سالانه میلیون‌ها نفر در تصادفات رانندگی در سراسر جهان، جان خود را از دست می‌دهند. جک ویست می‌گوید که رویکرد شرکت اینتل در مورد «ایمنی»، این است که از یک رویکرد قطعی استفاده کند؛ نه یک ماشین که بر اساس احتمال وقوع محاسبه‌شده توسط هوش مصنوعی تصمیم‌گیری می‌کند. او می‌گوید: «در این روش، تصمیماتی که هوش مصنوعی توصیه می‌کند، بررسی می‌شود. زمانی که الگوریتم چیزی را توصیه می‌کند که ممکن است برای ایمنی مناسب نباشد، ما یک روش قطعی داریم که بر اساس آن ایمنی توصیه‌شده را مورد بررسی قرار می‌دهیم.»

خودروهای خودران مطمئناً، مزایای بزرگی نسبت به انسان‌ها دارند. آنها قادرند چیزهای اطراف خود را سریع‌تر درک کنند و می‌توانند همزمان به تمام جهات نگاه کنند. زمان واکنش آنها نیز، سریع‌تر است و خسته نمی‌شوند.

پل نیومن می‌گوید: «تنها چیزی که ما به عنوان انسان در آن خوب عمل نمی‌کنیم، این

است که تمرکز خود را از دست می‌دهیم. حواس مان پرت می‌شود و این تنها اشتباهی است که یک ماشین انجام نمی‌دهد. چه تعداد تصادف به خاطر خستگی و چیزهای احمقانه رخ می‌دهد؟ علاوه بر این، خودروهای خودران می‌توانند از یکدیگر یاد بگیرند. سیستم‌های هوش مصنوعی اوکسبوتیکا (Oxbotica) که اساس عملکرد وسایل نقلیه آن را تشکیل می‌دهند، هر شب به روزرسانی می‌شوند. داده‌های جدید و سناریوهای جدیدی که وسایل نقلیه خودشان یافته‌اند. چه در دنیای واقعی و چه در شبیه‌سازی‌های مجازی. می‌توانند به دانش سیستم‌ها افزوده شوند. در یک روز، صدها ساعت تجربه و دانش جدید می‌توانند به درک وسیله نقلیه از نحوه رانندگی افزوده شوند. هر خودرویی می‌تواند تجربه خودروهای دیگر را داشته باشد و هنگامی که هر روز، هزاران یا میلیون‌ها وسیله نقلیه حرکت می‌کنند، این درک جمعی، به‌طور تصاعدی رشد می‌کند.»

فناوران معتقدند ماشین‌هایی از این دست در صورتی می‌توانند در دنیای واقعی کارکرد داشته باشند که به مردم در مورد چگونگی کارکرد واقعی و مزایای آنها توضیحاتی داده شود. نیومن اضافه می‌کند: «ما باید از این رویکرد که این موارد (نحوه کارکرد آنها) به‌نوعی عجیب و غریب و غیر قابل توضیح هستند فاصله بگیریم. آنها قطعات نرم‌افزاری مهندسی‌شده هستند؛ فرایندی پیرامون آنها وجود دارد و ما می‌توانیم بگوییم از آنها انتظار داریم چگونه رفتار کنند.»



بخش چهارم

خطرات هوش مصنوعی

بین سال‌های ۲۰۱۵ تا ۲۰۲۰، افراد متقاضی دریافت روادید برای ورود به انگلستان جهت کار، تحصیل یا دیدار با عزیزان‌شان، مدارک را به روش معمول پرمی کردند و سپس ارزیابی این داده‌ها توسط الگوریتمی انجام می‌شد. الگوریتم آنها را بر اساس سه رنگ قرمز، کهربایی یا سبز رتبه‌بندی می‌کرد. از میان کسانی که با رنگ سبز ارزیابی شدند، ۹۶/۳ درصد از آنها فوری تأیید شدند. افرادی که با رنگ قرمز مشخص شدند، جزو «خطرناک‌ترین» دسته بودند و به‌طور خودکار رد نشدند، اما تحت بررسی‌های بیشتری قرار گرفتند و کارمندان ارشد برای بررسی داده‌ها و تصمیم‌گیری نهایی وارد عمل شدند. این روند نیمه‌خودکار که توسط وزارت کشور اجرا شد، در نهایت ۴۸ درصد از متقاضیان با رنگ قرمز را تأیید کرد. کسانی که از این روند استفاده کردند به تصمیمات آن اعتماد داشتند.

اما مشکلی وجود داشت. اگرچه هدف به‌کارگیری این سیستم، قابل ستایش بود. کاغذبازی اداری کمتر، سریع‌تر و کارآمدتر کردن تقاضاهای دریافت روادید. اما فناوری زیربنایی آن، معروف به «ابزار پخش» (Streaming Tool) ناقص بود. این فناوری درگیر داده‌ها، شفافیت و مشکلات فرایندی شد که به تصمیم‌گیری‌های ناعادلانه منجر شد. به‌طور خاص، در مورد پرخطر بودن افراد، به جای بررسی دقیق معیارهای فردی، بر اساس کشور مبدأ افراد، قضاوت می‌کرد. «کوری کرایدر» (Cori Crider)، بنیان‌گذار گروه حقوقی «فاکس گلاو» (Foxglove) می‌گوید: «آنها فهرستی از کشورهای ناخوشایند داشتند که صرفاً با دریافت تقاضای روادید از یکی از این کشورها. البته از ارائه نام این کشورها به ما خودداری کردند. به احتمال زیاد آن شخص در دسته رنگ کهربایی یا قرمز، رتبه‌بندی می‌شد.»

در بهار ۲۰۲۰، فاکس گلاو، با مطرح کردن این بحث که این فرایند الگوریتمی، قوانین برابری و حفاظت از داده‌ها را نقض کرده، این فرایند را به چالش کشید و چند روز پیش از اینکه پرونده به دادگاه برسد، دولت انگلستان با اعتراف به اینکه «به‌طور کلی در ابزار پخش، مسائلی مانند سوگیری ناخودآگاه و استفاده از ملیت» نیازمند بررسی دقیق‌تر است، استفاده از این ابزار را کنار گذاشت. هرچند نقض هرگونه قانون توسط این سیستم رد شد.

چند ماه بعد، در میانه شیوع ویروس کرونا، یک الگوریتم مشابه تلاش کرد نمرات «آزمون سطح ای» (A-Level Exam) دانش آموزان در انگلستان را پیش بینی کند. این الگوریتم نیز، از داده های گذشته برای تصمیم گیری در مورد آینده افراد استفاده کرد. این الگوریتم نیز با مشکل مواجه شد؛ دانش آموزان اعتراضات گسترده ای انجام دادند و در خیابان ها شعار «مرگ بر الگوریتم» سر دادند.

هیچ کدام از این سیستم ها، نه الگوریتم های پیچیده خاصی را دربر می گرفتند و نه از هوش مصنوعی استفاده می کردند. اما آنچه هر دو نشان می دهند، ماهیت ذاتی مخاطره آمیز داده ها است. هرچقدر هم بخواهیم فرض را بر بی طرفی داده ها بگذاریم، باید گفت به ندرت بی طرف هستند، چون به هر حال توسط انسان ها ایجاد شده اند و تعصبات و پیش داوری های ما را به تصویر می کشند. سیستم صدور روادید، بر اساس داده هایی بنیان گذاشته شد که مفروضات مغرضانه ای راجع به کشورهای مبدا افراد داشت. سیستم امتحانات نیز نتایج را تا حدودی بر اساس سابقه مدارس شاگردان پیش بینی می کرد. بنابراین هر دو مورد، زنگ خطری را درباره آینده به صدا درمی آورند.

سوگیری

از آنجایی که داده ها در مرکز هوش مصنوعی قرار دارند، بنابراین نتیجه منطقی این است که هوش مصنوعی عاری از سوگیری ها نیست. داده های ورودی نامناسب در یک سیستم به داده های خروجی نامناسب منجر می شود. سه شکل متداول سوگیری وجود دارد. اگرچه تحقیقات بیش از ۲۰ نوع سوگیری همراه با دیگر تبعیض ها و بی عدالتی ها را شناسایی کرده. که می تواند در تنظیمات هوش مصنوعی تأثیرگذار باشد. یک الگوریتم با «سوگیری نهفته»، نتایجش را با ویژگی های داده ها از قبیل جنسیت، نژاد، درآمد و موارد دیگر مرتبط می کند و به این ترتیب، به طور بالقوه اشکال تاریخی سوگیری ها تداوم می یابد. به عنوان نمونه، یک سیستم شاید «یاد بگیرد» که پزشکان همگی مرد هستند؛ زیرا داده های تاریخی ای که این سیستم توسط آن آموزش دیده، پزشکان را مذکر نشان می دهد (شرکت آمازون مجبور بود ابزار استخدام هوش مصنوعی را که از آن استفاده می کرد، کنار بگذارد؛ زیرا این ابزار توسط داده های مربوط به ۱۰ سال سوابق شغلی،

آموزش دیده بود و از آنجایی که در طول این مدت، مردان بیشتر از زنان استخدام شده بودند، هوش مصنوعی فکر می‌کرد مردان بهتر از زنان هستند). به خاطر «سوگیری‌های شناختی»، نتایج توسط مجموعه داده‌ای که بیش از حد نمایانگر یک گروه خاص از افراد است، تحریف می‌شوند (هوش مصنوعی که برای قضاوت در یک مسابقه زیبایی مورد استفاده قرار گرفته بود، بیشتر سفیدپوستان را انتخاب کرد، چون مجموعه داده‌ای که توسط آن آموزش دیده بود، بیشتر حاوی تصاویر سفیدپوستان بود). به علت «سوگیری تعاملی»، یک سیستم، از پیش‌داوری‌های مردم هنگام تعامل با آن یاد می‌گیرد (چت‌بات مایکروسافت، به نام «تی» (Tay) که در سال ۲۰۱۶ در شبکه توییتر راه‌اندازی شد، در ۲۴ ساعت اول بهره‌برداری از آن انسجام بیشتری داشت، اما تمام الفاظ مربوط به تبعیض نژادی و جنسیتی را که مردم ارسال کرده بودند، تکرار کرد).

همه این سوگیری‌ها، خطرات استفاده از داده‌ها را برای پیش‌بینی نتایج آینده نشان می‌دهند. در هوش مصنوعی، این مشکل رسماً اعلام شده است. به عنوان مثال، اجرای قانون را در نظر بگیرید. به‌ویژه در ایالات متحده آمریکا، پلیس هنگامی که می‌خواهد بررسی کند که یک فرد خاص چقدر احتمال دارد دوباره تخلف کند یا آیا به او اجازه ارائه وثیقه داده شود، یا چه زمانی احتمال افزایش جرایم محتمل است، از سیستم‌های پیش‌بینی استفاده می‌کند. مشکل این است که چنین سیستم‌هایی به‌ناچار افراد و جوامعی را که در گذشته مورد توجه ویژه پلیس بوده‌اند، هدف قرار می‌دهند. «رنه کامینگز» (Renée Cummings)، پژوهشگر متخصص در هوش مصنوعی در حوزه عدالت کیفری و تأثیرات تبعیض آمیز آن، توضیح می‌دهد: «اگر از داده‌های مشکل‌ساز استفاده کنید کارکرد مأموران پلیس نیز دچار مشکل خواهد شد.» در ایالات متحده آمریکا، احتمال دستگیری سیاه‌پوستان دوبرابر بیشتر از سفیدپوستان بوده و احتمال کشته‌شدن مردان سیاه‌پوست، دوونیم برابر بیشتر از مردان سفیدپوست است. کامینگز می‌گوید: «این واقعیت از پیش‌داوری‌ها سرچشمه می‌گیرد. هوش مصنوعی واقعاً تعصبات و تبعیضی را که در این سیستم وجود داشته تقویت کرده است.»

تحقیق سال ۲۰۱۶ صورت گرفته توسط «پروپابلیکا» (ProPublica) نشان داد که چگونه نرم‌افزار پیش‌بینی به نام «کومپاس» (COMPAS)، در مورد سیاه‌پوستان سوگیری

داشت. مطالعه بیشتر روی این ابزار نشان داد که «این نرم افزار برای اسپانیایی زبان ها به خوبی تنظیم نشده بود». گویا این نرم افزار ریسک های بیشتری را برای اسپانیایی زبان ها پیش بینی می کرد و مشخص شد که توانایی انجام پیش بینی های دقیق را ندارد. هوش مصنوعی همچنین می تواند اعمال نادرست پلیس را تداوم بخشد. مطالعه محققان مؤسسه «AINow» دانشگاه نیویورک درباره ۱۳ حوزه قضایی ایالات متحده آمریکا که در آنها، ابزارهای پیش بینی پلیس در حال فعالیت بوده اند، نتیجه گیری کرد که «اقدامات غیرقانونی پلیس می تواند موجب تحریف داده های جمع آوری شده شود که به این ترتیب خطر استفاده از داده های کثیف در اجرای قانون و دیگر موارد ایجاد می شود».

حتی هنگامی که داده های مربوط به نژاد افراد از سیستم های پیش بینی هوش مصنوعی پلیس کنار گذاشته می شود (الزام مکرر به برابری در مقابل قانون)، مشکلات همچنان باقی است. «راشیدا ریچاردسون» (Rashida Richardson)، از نویسندگان مؤسسه «AINow» که در حال حاضر استاد میهمان در دانشکده حقوق «رانگروز» (Rutgers Law School) است، اشاره می کند که انواع مختلفی از داده ها وجود دارند که می توانند نشان دهنده نژاد باشند. مکان یکی از آنهاست. اگر در منطقه ای زندگی می کنید که به دلیل نژادها و قومیت های مختلف نسبت به دیگر مناطق بیشتر تحت نظارت پلیس قرار دارد، در نتیجه گزارش های پلیسی بیشتری وجود خواهد داشت و بنابراین احتمال اینکه هوش مصنوعی منطقه شما را جرم خیز تلقی کند بسیار زیاد است. سایر داده هایی که می توانند به تحریف منجر شوند شامل سن و ارتباطات اجتماعی هستند. یکی از ابزارهای پیش بینی پلیس در کشور هلند، از داده هایی شامل اطلاعات جمعیت شناختی، مانند تعداد خانوارهای تک والدی، تعداد افرادی که از مزایای اجتماعی استفاده می کنند و تعداد مهاجران از کشورهای به جز غرب، به عنوان عوامل تعیین احتمال وقوع جرم در منطقه ای خاص استفاده می کنند. مشکل همه اینها، این است که داده های گردآوری شده ممکن است گاهی به نکات مفیدی اشاره کنند، اما بی طرفانه پیش بینی نمی کنند. ریچاردسون توضیح می دهد: «داده های پلیس به احتمال زیاد بیشتر منعکس کننده اقدامات، سیاست ها و اولویت های اداره پلیس است تا روندهای جرم و جنایت.» وی این سؤال را مطرح می کند که آیا فناوری های پیش بینی

پلیس، بازتاب آنچه در جوامع اتفاق می افتد هستند یا خیر؛ «واقعیت این است که اگر سیستم پیش بینی بیشتر بر داده های پلیس متکی باشد، همواره نوعی نقص دائمی در آن وجود خواهد داشت که استفاده از آن را برای انجام هر کار پلیسی دشوار می کند.»

در حال حاضر شواهد کمی وجود دارد که نشان دهد ابزارهای ارزیابی خطر هوش مصنوعی یا ابزارهای پیش بینی، واقعاً مفید هستند. مطالعاتی که نشان می دهند ابزارهای تصمیم گیری الگوریتمی می توانند پیش بینی های بهتری نسبت به انسان ها داشته باشند، بسیار اندک هستند و از آنجایی که نیروهای پلیس عموماً درباره استفاده خودشان از فناوری های هوش مصنوعی مایل به صحبت نیستند، داده های آنها اغلب مستقلاً تحلیل و تأیید نشده است. با این حال، بررسی انجام شده توسط اداره پلیس لس آنجلس درباره سیستمی به نام «پردپل» (PredPol) که در چندین مکان در ایالات متحده آمریکا استفاده می شود، حاکی از آن است که نتیجه گیری در مورد اثربخشی این سیستم در کاهش جرایم مربوط به وسایل نقلیه یا جرایم دیگر، دشوار بود.

در ژوئن ۲۰۲۰، بیش از ۱۴۰۰ ریاضیدان، نامه ای سرگشاده را امضا کردند و اظهار داشتند که به عقیده آنها متخصصان این رشته نباید در زمینه این سیستم ها با پلیس همکاری کنند. همچنین خواستار آن شدند که بازرسانی برای نظارت بر سیستم هایی که از قبل وجود دارند معرفی شوند. محققان نوشتند: «به راحتی می توان این اشتباه را انجام داد و نژادپرستی را در ظاهر و پوشش علمی رواج داد.» «کوری کرایدر» که چالش های حقوقی برای سیستم های آماری ناقص در انگلستان ایجاد کرد، می افزاید اغلب به نظر می رسد این نوع فناوری ها ضد گروه هایی استفاده می شوند که شاید وسیله ای برای به چالش کشیدن آنها نداشته باشند. او می گوید: «ظاهراً بسیاری از سیستم ها صرفاً برای مدیریت انبوه افرادی که قدرت، سرمایه اجتماعی و پول بسیار کمتری دارند، مورد استفاده قرار می گیرد. من فکر می کنم که واقعاً روند نگران کننده ای وجود دارد. آن مدیریت الگوریتمی، راهی برای مهار و نظارت بر افراد فقیر رنگین پوست است.»

فقط بخش پلیس نیست که ممکن است با مشکلات پیش داور هوش مصنوعی مواجه باشد. مسکن، اشتغال و امور مالی، از جمله حوزه هایی هستند که متحمل مشکلات سوگیری هوش مصنوعی شده اند. در حال حاضر بهداشت و درمان جدیدترین

حوزه‌ای است که با این مشکلات دست‌به‌گریبان است. به‌عنوان مثال، در ایالات متحده آمریکا، پژوهشگران دانشگاه کالیفرنیا کشف کردند که الگوریتم استفاده‌شده توسط بیمه‌ها و بیمارستان‌ها برای مدیریت مراقبت بهداشتی حدود ۲۰۰ میلیون نفر در سال، به بیماران سیاه‌پوست رتبه خطر کمتری نسبت به بیماران سفیدپوست (با همان میزان بیماری) می‌داد. این موضوع به این دلیل بود که نمرات سلامتی افراد را تا حدی با توجه به هزینه‌ای که در یک سال برای مراقبت‌های درمانی صرف کرده‌اند، تعیین می‌کرد و فرض بر این است که افراد بیمار، بیشتر از افراد سالم هزینه می‌کنند.

داده‌ها همچنین نشان دادند که هزینه مراقبت‌های درمانی بیماران سیاه‌پوست، ۱۸۰۰ دلار کمتر از بیماران سفیدپوست با همان مشکلات مزمن سلامتی بود. نویسندگان این تحقیق نوشتند: «پول کمتری برای بیماران سیاه‌پوستی که همان میزان مراقبت درمانی نیاز دارند، هزینه می‌شود.» بنابراین، الگوریتم به اشتباه نتیجه می‌گیرد که بیماران سیاه‌پوست از بیماران سفیدپوست به همان اندازه بیمار، سالم‌تر هستند.» در نتیجه، افراد سیاه‌پوست برای درمان‌های تخصصی‌تر، کمتر ارجاع داده شدند. اگر شرایط مساوی می‌بود، باید ۴۶/۵ درصد از بیماران سیاه‌پوست ارجاع داده می‌شدند؛ در حالی که این رقم، ۱۷/۷ درصد بود.

محققان نگران هستند که هر چقدر هوش مصنوعی پیچیده‌تر شود، اتصالات ایجادشده توسط الگوریتم‌ها نیز مبهم‌تر شود. شناسایی متغیرها برای انواع خاصی از داده‌ها دشوارتر خواهد شد و ماشین‌ها شاید بین انواع خاصی از اطلاعات پیوند برقرار کنند که انسان‌ها آنها را به همدیگر مرتبط نمی‌کنند یا به دلیل مقیاس اطلاعات پردازش‌شده نمی‌توانند آنها را ببینند. اولین نشانه‌های این رخداد، در حال حاضر مشهود است. برای نمونه، در مارس ۲۰۱۹، وزارت مسکن و شهرسازی ایالات متحده آمریکا (HUD) فیس‌بوک را به تبعیض در تبلیغات هدفمندش برای مسکن متهم کرد. قوانین تبلیغات در ایالات متحده آمریکا، تبعیض علیه مردم بر اساس رنگ پوست، نژاد، کشور زادگاه افراد، جنسیت، دین، ناتوانی و وضعیت خانوادگی‌شان را منع کرده است. در حالی که سیستم فیس‌بوک مستقیماً مسبب چنین تبعیضی نبود، اما معلوم شد که علایق بیان‌شده افراد به‌صورت آنلاین به محرومیت‌شان از برخی تبلیغات منجر

شده است.

«ساندرا واچر» (Sandra Watcher)، دانشیار و محقق ارشد حقوق و اخلاق هوش مصنوعی در مؤسسه اینترنت آکسفورد می گوید: «یک الگوریتم افراد را لزوماً بر اساس گرایش جنسی یا رنگ پوست شان گروه بندی نمی کند.» وی می گوید الگوریتم ها افراد را با توجه به رفتارهای مشابه شان گروه بندی می کنند؛ «این شباهت ها می تواند این باشد که همه آنها کفش سبز دارند، یا شباهت ها می تواند این باشد که آنها در یک سه شنبه، غذای هندی می خورند یا اینکه بازی های ویدئویی انجام می دهند.» مسئله این است که چنین ارتباطاتی می توانند به استنباط های کاملاً نادرستی منجر شوند. برای مثال، ممکن است یک الگوریتم نتیجه بگیرد که افرادی که کفش سبز می پوشند تمایل دارند یا تمایل ندارند که وام های خود را به موقع بازپرداخت کنند. چنین نتیجه گیری ای از نظر هر انسانی شاید غیرمنتظره باشد، اما از نظر هوش مصنوعی کاملاً منطقی به نظر می رسد.

واچر می نویسد حتی اگر بانکی بتواند توضیح دهد که کدام داده ها و متغیرها برای تصمیم گیری استفاده شده اند (به عنوان مثال، سوابق بانکی، درآمد و کد پستی)، این تصمیم گیری، به استنباط از این منابع وابسته است. برای نمونه، متقاضی، وام گیرنده قابل اعتمادی نیست؛ «این، یک فرض یا پیش بینی در مورد رفتار آینده است که در زمان تصمیم گیری نمی تواند تأیید یا رد شود.» او این بحث را مطرح می کند که مقررات حفاظت از داده باید برای رسیدگی به مسائلی که در نتیجه استنتاج نادرست ماشین ها درباره ما به وجود می آیند، تکامل یابند. سازمان هایی که از چنین سیستم هایی استفاده می کنند باید ثابت کنند ارتباطاتی که ایجاد می کنند معقول هستند. آنها باید بگویند چرا برخی داده ها مرتبط هستند، چرا استنباط ها برای تصمیماتی که گرفته می شود مهم هستند و آیا این فرایند، دقیق و قابل اعتماد است. به گفته واچر، در غیر این صورت افراد بیشتری به خاطر هوش مصنوعی که تصمیمات ناعادلانه می گیرد رنج خواهند کشید.

نظارت

اگر هوش مصنوعی هنگام پردازش داده ها مشکل ساز باشد، هنگام پردازش تصاویر

مربوط به انسان‌ها نیز. به‌ویژه از طریق فناوری تشخیص چهره. کاستی‌های شدیدی را نشان می‌دهد.

نحوه کار این سیستم‌ها، دست‌کم در تئوری نسبتاً ساده است. اگر یک دوربین نظارتی متصل به سیستم هوش مصنوعی چهره‌ای را شناسایی کند، اندازه‌گیری‌هایی بین اجزای صورت مانند فاصله بین پیشانی و چانه انجام و نقشه بیومتریکی ایجاد می‌شود. از این نقشه برای ایجاد «اثر چهره» (که در واقع، یکسری اعداد است) استفاده می‌شود. سپس این مورد با سایر «اثرهای چهره» موجود در پایگاه داده برای یافتن یک تطبیق احتمالی مقایسه می‌شود. الگوریتم‌ها برای تطبیق چهره توسط مجموعه داده‌ها آموزش داده شده‌اند.

نیروهای پلیس در سراسر جهان از چنین هوش مصنوعی برای شناسایی افرادی که قبلاً به‌عنوان مجرم یا مجرم بالقوه وارد سیستم‌هایشان شده‌اند، استفاده می‌کنند. اگر هنگام استفاده تطبیق چهره به‌صورت زنده انجام شود به افسران حاضر در موقعیت هشدار داده می‌شود و می‌توانند مظنون را ردیابی کنند. برخی دوربین‌های دارای سیستم‌های شناسایی چهره به‌صورت زنده، موقعیت را به‌وضوح نشان می‌دهند، اما اغلب مدل‌های پیشرفته‌تر، بسیار شبیه دوربین‌های مداربسته CCTV هستند. یکسری از دوربین‌های پیشرفته در سال ۲۰۲۰ راه‌اندازی شده که از شش الگوریتم هوش مصنوعی استفاده می‌کنند؛ تصاویری با رزولوشن ۱۲ مگاپیکسل ارائه می‌دهند و تشخیص چهره را با شمارش جمعیت و خواندن پلاک خودرو ترکیب می‌کنند. این سیستم می‌تواند محیط و صف‌ها را تحت نظارت قرار دهد و حتی می‌تواند تشخیص دهد که افراد از کلاه‌های ایمنی استفاده می‌کنند یا نه.

حداقل، در تئوری که به این شکل است. در عمل، دقت بالاتر از آزمون‌های آزمایشگاهی، اغلب در دنیای واقعی تکرار نمی‌شوند. نور کم، همانند فیلم‌های ویدئویی بی‌کیفیت، می‌تواند فرایند شناسایی چهره را کاملاً مختل کند. خورشید (به‌عنوان منبع نور) نیز می‌تواند دردسر بزرگی باشد. در یک آزمایش اولیه که روی سیستم‌های «شناسایی چهره به‌صورت زنده»، توسط پلیس انگلستان صورت گرفت، ارزیابی شد که ۲/۴۷۰ چهره با آنچه در پایگاه داده پلیس نگهداری می‌شود مطابقت دارد، اما در واقعیت فقط ۸

درصد موارد مطابقت داشتند و موارد دیگر که به اصطلاح با همدیگر مطابقت داشتند، نادرست بودند.

دقت، مشکل ساز است. از آنجایی که بسیاری از سیستم‌ها عمدتاً از طریق تصاویر افراد سفیدپوست آموزش دیده‌اند، هنگام تشخیص چهره افراد سیاه‌پوست دقت کمتری دارند؛ بنابراین شناسایی‌های نادرستی انجام می‌دهند. در ژانویه ۲۰۲۰، مرد سیاه‌پوست ۴۱ ساله اهل میشیگان به نام «رابرت ویلیامز»، به دلیل اشتباه الگوریتم تشخیص چهره، ۳۰ ساعت را در بازداشت به سر بُرد. سیاه‌پوست دیگری، «مایکل اولیور» (Michael Oliver) ۲۶ ساله از دیترویت نیز به دلیل شناسایی چهره اشتباه، از پلیس شکایت کرد.

مسائل مربوط به حریم خصوصی نیز مطرح هستند. دوربین‌های مدار بسته اکنون همه‌جا فراگیر شده‌اند و ما در مورد محل قرارگیری و نحوه استفاده از تصاویرشان حق اظهار نظر و مشارکت در تصمیم‌گیری نداریم. البته به مرور می‌توان آنها را جهت استفاده برای تشخیص چهره ارتقا داد. نیروهای پلیس ادعا می‌کنند که تشخیص چهره می‌تواند زندگی مان را ایمن‌تر کند و اینکه آنها ابزار مهمی در مبارزه با جرایم هستند. اما در کنار چنین مزایایی، باید ملاحظات مربوط به آزادی مدنی نیز لحاظ شود. در انگلستان، چندین نفر به دلیل اینکه صورت‌شان را پوشاندند تا چهره‌شان توسط این فناوری اسکن نشود، بازداشت و مورد بازجویی قرار گرفتند. حتی نوجوانی پس از اینکه توسط دادگاه آزاد شد، مجدداً به اشتباه بازداشت شد (در این مورد، می‌توان تقصیر را گردن پایگاه داده قدیمی انداخت). یکی از خطرات این است که چون در حال حاضر نظارت، بسیار گسترده و فراگیر است، نظارت کردن بسیار عادی شده است. بسیاری از افراد با دوربین‌های امنیتی خانگی مجهز به قابلیت تشخیص حرکت آشنا هستند. «رینگ» (Ring) متعلق به آمازون، محبوب‌ترین آنهاست. بنابراین ممکن است نظارت را به عنوان یک مزیت در زندگی روزمره در نظر بگیرند. چنین افرادی تمایل دارند در محله‌های مرفه‌تری زندگی کنند که احتمال کمتری در شناسایی اشتباه صاحبان خانه توسط این فناوری وجود دارد. در نتیجه این افراد از خطرات احتمالی این فناوری چندان آگاهی ندارند، اما به این معنا نیست که معایبی وجود ندارد.

شاید جنبه‌های مثبت و منفی دنیای نظارت در سطح ملی بیشتر آشکار شود. قطعاً در سطح ملی میزان واقعی قدرت آن، کاملاً مشخص است. کشور هند در هفته‌های ابتدایی شیوع ویروس کرونا در آوریل ۲۰۲۰، از هوش مصنوعی استفاده کرد. با گسترش موج اول؛ نیروهای پلیس در دهلی و بمبئی، برای اجرای قرنطینه ملی و بهبود فاصله‌گذاری اجتماعی، هواپیماهای بدون سرنشین را برای نظارت بر مردم به پرواز درآوردند (یک استارتاپ محلی در بمبئی با اهدای ۴۵ پهپاد کمک کرد). دوربین‌های موجود، تصاویری از اجتماعات غیرقانونی را ضبط کرده و بلندگوها پیام‌هایی را با صدای بلند اعلام کردند که به گروه‌ها دستور می‌داد، پراکنده شوند. پلیس که در همان لحظه در حال مشاهده فیلم بود، توانست به سرعت افسران پلیس را به مکان مورد نظر اعزام کند.

در این میان، در «آمریتسار» (Amritsar)، دومین شهر بزرگ ایالت پنجاب، نظارت هوایی سریع‌تر و قوی‌تر شد. آزمایشگاه‌های استارتاپ «اسکای لارک» (Skylark)، هواپیماهای بدون سرنشینی را به پرواز درآوردند که از «بینایی رایانه» جهت اندازه‌گیری میزان نزدیکی افراد به یکدیگر استفاده می‌کردند (اگر افراد به اندازه کافی از یکدیگر فاصله داشتند، کادر سبزرنگی اطراف آنها ظاهر می‌شد و اگر کادر قرمز رنگی ظاهر می‌شد، به آن معنا بود که افراد خیلی به یکدیگر نزدیک هستند). «آمارجوت سینگ» (Amarjot Singh)، مدیرعامل اسکای لارک توضیح می‌دهد: «یک الگوریتم نقشه‌برداری وجود دارد که نقشه کل شهر را به ۱۰ بخش تقسیم می‌کند و ما فقط به پلیس می‌گوییم که پهپادها چه بخش‌هایی را می‌خواهند برای پرواز انتخاب کنند و پهپادها به‌طور خودکار در آن بخش به پرواز درمی‌آیند.» وی اضافه می‌کند: «اگر انسانی شناسایی شود، این انسان را علامت‌گذاری می‌کند و پیامی را به تلفن همراه افسران پلیسی که به آن منطقه نزدیک هستند، ارسال می‌کند.» این سیستم بعداً در شش شهر دیگر هم راه‌اندازی شد و در نهایت، ۲۰۰ جریمه از این طریق صادر و برخی خودروها توقیف شدند. البته استفاده از این فناوری، نگرانی‌های جدیدی درباره کنترل دولت بر شهروندان ایجاد می‌کند و رفتار مردم بی‌سروصدا از بالا مشاهده می‌شود.

هند تنها کشوری نیست که در آن چنین نظارت ناخوانده و سرزده‌ای توسط دولت‌های ایالتی و محلی در حال انجام شدن است. «استیون فلدشتاین»

(Steven Feldstein)، یکی از اعضای ارشد «برنامه دموکراسی، درگیری و حکومت» شرکت کارنگی می گوید: «کشورهای بیشتری نسبت به آنچه من انتظار داشتم، در حال تهیه، تدارک، خرید و سرمایه گذاری در این سیستم‌ها هستند.» وی اولین بررسی درباره نحوه استفاده از نظارت هوش مصنوعی، سیستم‌های شهر هوشمند، شناسایی چهره و راه اندازی پلیس هوشمند را در سراسر جهان انجام داد و تخمین زد که تا تابستان ۲۰۲۰، ۷۷ کشور از ۱۷۹ کشور جهان در حال استفاده از چنین فناوری‌هایی بودند. از دموکراسی‌های لیبرال، مانند انگلیس و ایالات متحده آمریکا تا کشورهای دیگری مانند روسیه و عربستان سعودی جزو این کشورها هستند. فلدشتاین به رابطه میان «سرمایه گذاری در نظارت» و «کشورهای دارای بودجه نظامی بالا» اشاره کرد.

این فناوری، از نظر دسترسی و جاه طلبی‌هایی که نسبت به آن وجود دارد، در حال رشد است. اکنون می‌توان افراد را از طریق عنبیه چشم، خالکوبی، نحوه راه رفتن‌شان (فناوری تشخیص نحوه راه رفتن) و بسیاری موارد دیگر شناسایی کرد. به زودی جایی برای پنهان شدن وجود نخواهد داشت. پژوهشگران در چین و انگلیس، سیستمی ساخته‌اند که می‌تواند افراد را طی چند هفته شناسایی کند، حتی اگر لباس و مدل موهای خود را تغییر داده باشند. همچنین سیستم‌هایی در حال ساخت است که ادعا می‌شود احساسات یا هیجانات مردم را بر اساس حالات چهره آنها تشخیص می‌دهند (کارشناسان می‌گویند چنین جاه طلبی‌هایی مشکل دار است و قبل از هر چیز باید دید که آیا این فناوری اصلاً باید ایجاد شود یا خیر). اغلب این فناوری‌ها، هنوز در مرحله آزمایشی قرار دارند، اما تصویر خوبی از مسیر کلی فناوری ارائه می‌دهند. برای چنین فناوری‌هایی تعداد زیادی داده‌های آموزشی، سخت افزارهای پیچیده و الگوریتم‌های یادگیری ماشین بهبود یافته تهیه می‌شود.

برخی مقامات با چنین پیشرفت‌هایی محتاطانه برخورد می‌کنند. در ایالات متحده آمریکا استفاده از سیستم‌های شناسایی چهره در برخی ایالت‌ها ممنوع شده است. در شهر پورتلند در ایالت اورگان، کامل‌ترین ممنوعیت اتخاذ شده است (حتی در آنجا کسب و کارها هم مجاز به استفاده از این فناوری نیستند). گزارش شده که اتحادیه اروپا نیز در حال در نظر گرفتن محدودیت‌های مشابهی است. از سویی دیگر، در کشورهایی مانند

هند، روز به روز بیشتر از چنین فناوری‌هایی استفاده می‌شود. مقامات رسمی در آنجا به خود می‌بالیدند که از این فناوری در شناسایی ۱۱۰۰ نفر درگیر در شورش‌های خشونت‌آمیز دهلی طی فوریه ۲۰۲۰ استفاده کردند. گزارش شده که «آمیت شاه» (Amit Shah)، وزیر کشور هند، در مورد افراد بی‌گناهی که توسط دوربین‌ها شناسایی شدند، گفت: «این، یک نرم‌افزار است. ایمان افراد را نمی‌بیند. لباس افراد را نمی‌بیند. فقط صورت آنها را می‌بیند و افراد از طریق صورت، گیر می‌افتند.»

چین بیشترین آزمایش‌ها را در مورد این فناوری انجام داده است. هیچ ملتی به اندازه این ابردولت کمونیست در پیگیری سیستم‌های نظارتی مشتاق نبوده است. به‌طور معمول اطلاعات عظیمی از شهروندان خود را جمع‌آوری می‌کند و فناوری شناسایی چهره را به‌طور گسترده‌ای اجرا کرده است. استیون فلدشتاین می‌گوید: «مدل چین، بسیار متمایز است. هیچ کشور دیگری از نظر تقلید از این رویکرد یا حتی داشتن آرزوی انجام نظارت کامل، واقعاً به گرد پای آن نمی‌رسد.» این کشور همچنین بزرگ‌ترین صادرکننده فناوری نظارت هوش مصنوعی است. فلدشتاین تخمین می‌زند که «هواوی» (Huawei)، «هایک‌ویژن» (Hikvision)، «داهوا» (Dahua) و «زدتی‌ای» (ZTE)، تجهیزات فناوری ۶۳ کشور را تأمین می‌کنند. هواوی، به‌تنهایی تجهیزات ۵۰ کشور را تأمین می‌کند که نسبت به شرکت غیرچینی بعدی، یعنی «ان‌ای‌سی» (NEC) ژاپن، با فاصله بسیار زیادی در صدر قرار دارد (ان‌ای‌سی در ۱۴ کشور جهان استفاده می‌شود).

یک مطالعه انجام‌شده توسط پژوهشگران مؤسسه تحقیقاتی شانگهای که در ژوئن ۲۰۲۰ منتشر شد، نشان می‌دهد که اگر این فناوری کنترل نشود، آینده ویران‌شهری (دیس‌توپیا) چگونه خواهد بود. این تحقیق، از یک شبکه دوربین‌های امنیتی، به همراه حسگرهای دیگر (که حرکت، دما و موارد دیگری را اندازه‌گیری می‌کرد) در یکصد هزار آسانسور در سراسر چین استفاده کرد و داده‌هایی را برای چندین الگوریتم فراهم کرد تا به تعیین اینکه چه زمانی رفتار افراد، «عادی» نیست، کمک کند. هدف این بود که رفتارهای بالقوه ناامن شناسایی شود (محققان گزارش می‌دهند که حتی فردی به آنها پول داد تا بتواند از این فناوری برای شناسایی هرگونه ازدحام غیرقانونی در برخی آپارتمان‌ها، استفاده کند). البته پژوهشگران امیدوار بودند که جنسیت، سن، ظاهر و شغل افراد را نیز

تعیین کند. رفتارهای بالقوه ناامن، تنها شامل «محل زندگی بیش از حد شلوغ» نمی‌شد؛ بلکه «تجارت مواد مخدر، شرکت‌های هرمی، روسپی‌گری و غیره» را نیز دربر می‌گرفت. پژوهشگران در مقاله خود نوشتند: «هر بار که آسانسور در یک طبقه متوقف می‌شد، عکسی گرفته می‌شد، صد هزار آسانسور توانستند برای هر طبقه، حدود یک میلیون رکورد ثبت کنند.» در مجموع، ۶۴۳ مورد شناسایی شد که به‌طور بالقوه غیرطبیعی رفتار می‌کردند. از این تعداد، بیش از یک‌صد مورد، اشتباه. به دلیل خطاهای موجود در تجهیزات یا نتیجه تصاویر جمع‌آوری شده از دوربین‌های امنیتی موجود در شبکه. بودند که در بلوک‌های آپارتمانی قرار نداشتند (مانند دوربین‌های موجود در آسانسورهای مراکز خرید یا ساختمان‌های اداری). در ۲۸۹ مورد، محققان متوجه موارد مشکوکی شدند و به این نتیجه رسیدند که مدیر ساختمان باید تلاش کند جزئیات بیشتری به دست آورد. بر اساس این تحقیق معلوم شد سه مورد از موارد علامت‌گذاری شده، افرادی هستند که خدمات تهیه غذا از خانه‌های خود ارائه می‌دادند و شش مورد به ازدحام بیش از حد اشاره می‌کردند.

با وجود مقیاسی که این سیستم در آن کار می‌کرد، سیستمی نسبتاً ابتدایی بود و در لحظه کار نمی‌کرد و اندازه زیاد داده‌های گردآوری شده، به این معنی بود که محققان، سن و مشاغل احتمالی افراد را تجزیه و تحلیل نکردند. اما پتانسیل این کار وجود داشت و اینکه این پتانسیل نه فقط برای تشخیص کارهای اشتباه، بلکه برای کنترل و تجسس در زندگی خصوصی افرادی که هیچ کار اشتباهی مرتکب نشده‌اند. پیامدهای نگران‌کننده‌ای دارد و ممکن است بر جنبه‌های مختلف زندگی روزمره افراد، از کارهای روزمره گرفته تا میزان بهره‌وری افراد در محل کار و تعیین اهداف برای آنها و زیر نظر داشتن ما در فروشگاه‌ها و فرودگاه‌ها نظارت کند.

یک نمونه از سرکوب ناشی از نظارت جمعی که منتقدان دولت چین به آن اشاره می‌کنند، به گروه اقلیت اکثرأ مسلمان این کشور یعنی «اویغورها» (Uighurs) مربوط است. نیویورک تایمز گزارش داد که در یک منطقه، فقط برای یک ماه در طول سال ۲۰۱۹، ۵۰۰ هزار بار از تشخیص چهره استفاده شد تا مشخص شود آیا افرادی که از روبه‌روی دوربین گذشته‌اند، اویغور هستند یا خیر. این روزنامه به نقل از یک شرکت گزارش داد:

«اگر در ابتدا یک اویغور در محله‌ای زندگی کند و ظرف ۲۰ روز شش اویغور ظاهر شوند، این سیستم بلافاصله هشدارهایی را برای پلیس محلی ارسال می‌کند.» چنین توانایی فناورانه‌ای در آینده افزایش خواهد یافت.

جعل عمیق (Deepfake)

در اواخر سال ۲۰۱۷، یک کاربر ناشناس در «ردیت» (Reddit)، دنیای آنلاین را تغییر داد. با استفاده از «جعل عمیق»، شروع به اشتراک گذاری فیلم‌های پورنوگرافی از افراد مشهور کرد و از طریق یادگیری عمیق، چهره‌های افراد مشهور را جایگزین چهره‌های اصلی کرد. تیلور سوئیفت، گال گدوت، میسی ویلیامز و اسکارلت جو هانسون، برخی از قربانیان این حادثه بودند. از آن زمان به بعد، جعل عمیق گسترده شد.

البته فتوشاپ، چندین دهه است که وجود دارد، اما جعل عمیق بسیار پیچیده‌تر است و می‌تواند به میزان زیادی قانع‌کننده باشد. برخلاف فتوشاپ، آنها به جای صرفاً وابستگی به عامل انسانی برای قرار دادن تصاویر از منابع مختلف در کنار یکدیگر، تصاویر اصلی را از طریق یک برنامه هوش مصنوعی به دست می‌آورند که قادر است از داده‌های ارائه‌شده برای تولید داده‌های جدید استفاده کند. «شبکه‌های مولد تخاصمی» (GAN-Generative Adversarial Network) که در سال ۲۰۱۴ ایجاد شدند، رایج‌ترین ابزار مورد استفاده در این حوزه هستند. هر چقدر مجموعه داده‌های به کار گرفته‌شده بزرگ‌تر باشد، تصاویر ایجادشده واقعی‌ترند.

جعل عمیق، یکی از مصادیق چالش‌هایی است که هوش مصنوعی برای جامعه ایجاد کرده است. استفاده از آن می‌تواند کاملاً بی‌ضرر باشد، اما همچنین می‌تواند سبب آسیب واقعی شود. فناوری پشت آن، از اولین تصاویر بی‌کیفیت ایجادشده در ردیت، تا به امروز بسیار پیشرفت کرده است. این فناوری به‌طور تصادفی در حال رشد است و در این فرایند از طریق آموزش‌های آنلاین رایگان و کد منبع‌باز در حال دموکراتیک (مردمی) شدن است. در جولای ۲۰۱۹، شرکت امنیت سایبری تشخیص جعل عمیق به نام «سنسیتی» (Sensity)، ۱۴/۶۷۸ جعل عمیق را به‌صورت آنلاین پیدا کرد. این عدد در ژوئن ۲۰۲۰، به ۴۹،۰۸۱ مورد افزایش یافت. اغلب آنها پورنوگرافی

هستند و بیشتر زنان را هدف قرار می‌دهند. ارقام مربوط به سنسیتی در ژوئن ۲۰۲۰ آشکار کرد که هزار فیلم پورنوگرافی ایجاد شده توسط جعل عمیق، ماهیانه در سایت‌های اصلی ویدئویی بزرگسالان بارگذاری و میلیون‌ها بار دیده می‌شوند. آنچه مسئله را وخیم‌تر می‌کند، این است که هیچ قانونی برای رسیدگی به این مشکل وجود ندارد. قربانیان مجبور هستند صرفاً از قوانین موجود در زمینه کپی‌رایت، افترا و حقوق بشر استفاده کنند تا جعل‌های عمیق ایجاد شده خود را از اینترنت حذف کنند.

این مشکل فقط برای ستارگان سینما و خوانندگان نیست. کاربران تولیدکننده محتوا در یوتیوب، ستارگان اینستاگرام و زنانی که در «توییچ» (Twitch) به صورت زنده تصویرشان پخش می‌شود نیز هدف این فناوری قرار گرفته‌اند. در اکتبر ۲۰۲۰ مشخص شد که یک روبات هوش مصنوعی در پلتفرم پیام‌رسان تلگرام از فناوری جعل عمیق استفاده می‌کند تا لباس‌ها را از عکس‌های زنان حذف کند و اعضای برهنه بدن را به جای قسمت‌های پوشیده شده نشان دهد. در این راستا بیش از یکصد هزار زن قربانی شدند. بیشتر کسانی که این تصاویر را ایجاد می‌کنند می‌گویند از این فناوری علیه زنانی استفاده کردند که آنها می‌شناختند.

واقعیت این است که تعداد سرهای دیو هفت‌سر جعل آنلاین در حال بیشتر شدن است. به عنوان مثال، در تابستان ۲۰۱۹، روزنامه‌نگاران «آسوشیتدپرس» یک حساب جعلی «لینکدین» مرتبط با شخصیت‌های برجسته تأثیرگذار ایالات متحده آمریکا پیدا کردند که از تصویر پروفایل ساخته شده توسط شبکه مولد تخصصی استفاده می‌کرد. «هنری آژدر» (Henry Ajder)، در زمان صحبت با یک محقق «هوش تهدید» در شرکت سنسیتی، توضیح می‌دهد: «ما در لینکدین و توییتر مکرراً با آنها روبه‌رو می‌شویم، حتی در جست‌وجوی آنها نیستیم، با این موارد به‌طور اتفاقی مواجه می‌شویم. این تصاویر به‌راحتی در دسترس هستند، دسترسی به آنها مثل آب خوردن است.» در دسامبر ۲۰۱۹، فیس‌بوک صفحاتی را که از گرجستان و ویتنام نشئت می‌گرفتند، حذف کرد که شرکت سنسیتی تأیید کرد از تصاویر جعلی به عنوان عکس‌های پروفایل استفاده کرده بودند. در آگوست ۲۰۲۰، فیس‌بوک خودش اعلام کرد که ۱۳ پروفایل را حذف کرده که بر این باور بود منشاء آنها در روسیه بوده و به نظر می‌رسید که احتمالاً از

شبکه مولد تخصصی برای ایجاد تصاویر خود استفاده کرده است. فرایند تولید چنین پروفایل‌های جعلی، به مرور آسان‌تر می‌شود. نرم‌افزار تولید چهره منبع‌باز StyleGAN و StyleGAN2 شرکت «انویدیا» که از شبکه‌های مولد تخصصی استفاده می‌کند برای همه رایگان است و به راحتی برای ایجاد تصاویر افراد خیالی استفاده می‌شود.

به همین شکل صدا نیز در حال دستکاری شدن است و می‌تواند به کلاهبرداری‌های مالی کمک کند. شرکت امنیت سایبری «نیسوس» (NISOS) از ژوئن ۲۰۲۰، پرونده‌ای را (همراه با شواهد صوتی) مستند کرد که کارمندی در شرکت فناوری پیامی صوتی را دریافت می‌کند که ظاهراً از طرف مدیرعامل‌شان، پرداخت فوری پولی را درخواست کرده بود. خوشبختانه، فایل جعلی به اندازه کافی خوب نبود تا بتواند کارمند را فریب دهد، چراکه صدای جعلی به اندازه کافی به صدای مدیرعامل شبیه نبود. تجزیه و تحلیل نیسوس متعاقباً آشکار کرد که صدا دارای کیفیت بسیار متغیری بود و نت‌های بالا اغلب به اوج می‌رسیدند و اینکه سروصدای پس‌زمینه که آدم ممکن است در یک تماس عادی انتظار داشته باشد، وجود نداشت. این اولین موردی است که به‌طور کامل شرح داده شده و تلاش‌های آینده در زمینه جعل، بدون شک باورپذیرتر خواهند بود. در مورد متن تولیدشده توسط هوش مصنوعی، یک دانش آموز کالج آمریکایی به نام «لیام پور» (Liam Porr)، از «GPT-3» برای ایجاد پست‌های وبلاگی استفاده کرد که طی دو هفته انتشار آنلاین‌شان، ۲۶ هزار بار خوانده شدند. یکی از پست‌ها حتی در بالای سایت گردآورنده اخبار و فروم «اخبار هکر» (HackerNews) قرار گرفت. «پور» گفت که فقط چند نفر از افرادی که از GPT-3 مطلع بودند. این پرسش را مطرح کردند که آیا متنی که می‌خوانند توسط یک انسان نوشته شده است؟

در حال حاضر کیفیت «برنامه‌های ساختگی» به‌طور قابل توجهی متفاوت است. با وجود این، این برنامه‌های ساختگی می‌توانند برای بسیاری از افراد متقاعدکننده باشند، به‌ویژه در میان افرادی که از نظر اجتماعی یا سیاسی در طرفداری از پیامی خاص تعصب دارند. ویدئویی از «نانسی پلوسی»، رئیس مجلس نمایندگان ایالات متحده آمریکا که عمده‌برای مست‌نشان دادن او، سرعتش را کاهش داده بودند، میلیون‌ها بار پس از انتشارش در ماه می ۲۰۱۹ مشاهده شد. در کشور گابُن در اوایل سال ۲۰۱۹، ویدئویی از

رئیس جمهور این کشور. که برخی آن را جعل عمیق خواندند، اگرچه هیچ پژوهشی قادر نبود به طور قطعی این موضوع را اثبات کند. به ناآرامی های نظامی منجر شد. هنگامی که برنامه های ساختگی، قوی تر می شوند، ناگزیر تشخیص آن نیز دشوارتر خواهد بود. ابزارهای نرم افزاری، دستکاری متقاعدکننده فیلم، متن، صدا و تصاویر را برای هر کسی آسان خواهند کرد. پژوهشگران در تلاش اند ابزارهای هوش مصنوعی را برای شناسایی جعل های عمیق توسعه دهند، اما راهی طولانی در پیش دارند. به نظر می رسد که برای محافظت بهتر از افراد، قوانین باید به روز شوند. پلتفرم های شبکه های اجتماعی نیز به اجرای سیاست های قوی نیاز خواهند داشت. همچنین کسانی که قربانی شبکه های ساختگی می شوند، قادر خواهند بود به طور کامل از حقوق قانونی خود برای مطالبه غرامت استفاده کنند.

شکی نیست که شبکه های ساختگی قدرتمندتر، در راه هستند. بر اساس گزارشی از «سامسونگ نکست»، بازوی سرمایه گذاری غول فناوری کره جنوبی، در نوامبر ۲۰۲۰ بیش از ۱۷۰ شرکت استارت آپی در زمینه برنامه های ساختگی هوش مصنوعی در حال فعالیت بودند. این استارت آپ ها، هم شامل شرکت هایی مانند شرکت سنسیتی هستند که برای کاهش تهدید برنامه های ساختگی تلاش می کنند و هم شرکت هایی که در حال کار روی تولید محتوای جدید جهت استفاده در فیلم ها، صوت های ضبط شده، تبلیغات، اپلیکیشن ها و موارد دیگر هستند. برخی در تلاش برای ساخت سیستم های هوش مصنوعی هستند تا بتوانند همانند انسان، با مکث و بیان کامل و با صدای بلند بخوانند. برخی شرکت های دیگر به دنبال تبدیل مقاله های خبری نوشته شده به ویدئوها به میزبانی مجریان خبری هوش مصنوعی، قرار دادن چهره افراد در صحنه های فیلم مورد علاقه شان و ایجاد آواتارهای واقع گرایانه هستند که می توانند در صنعت بازی های چند میلیارد دلاری مورد استفاده قرار گیرند.

هوش مصنوعی در هالیوود نیز دارد راه خود را باز می کند. صنعت سینما برای دهه ها در حال استفاده از «تصاویر تولید شده توسط رایانه» (CGI) و جلوه های ویژه بوده است. ویدئوی ساختگی، این امکان را برای استودیوها فراهم می کند که بتوانند هزینه های فیلم برداری را کاهش دهند و از بازیگران به روش های جدید استفاده کنند. هنوز راه طولانی

در پیش است، اما «ویکتور ریپاربلی» (Victor Riparbelli)، مدیرعامل شرکت ویدئوی ساختگی به نام «سینتسیا» (Synthesia)، این طور پیش بینی می کند: «من فکر می کنم، ظرف مدت ۱۰ سال، می توانید فیلمی هالیوودی را با لپ تاپ تولید کنید.»

چنین پیشرفت های بالقوه ای یادآوری می کنند که همانند تمام فناوری ها به طور کلی و هوش مصنوعی به طور خاص، رسانه ساختگی نویدهای عظیم و همچنین چالش ها و تهدیدهایی را به همراه دارد. حتی در حال حاضر، برخی از روش های استفاده از آن نشان می دهد که توانایی بسیار بیشتری از فیلم های جعلی و پیام های کذب دارد. در فوریه ۲۰۲۰، در کمپین انتخاباتی هند در دهلی، «مانوج تیواری» (Manoj Tiwari)، رئیس حزب «بهاراتیا جاناتا» (BJP)، با انتشار دو فیلم توانست یک جمعیت حدود ۱۵ میلیون نفری را مخاطب خود قرار دهد. در یکی از فیلم ها او به زبان انگلیسی صحبت کرد، در فیلم دیگر، با کمک یک دوبلور و الگوریتم همگام سازی لب (لیپ سینک)، حالت دهان تیواری تغییر کرد و به نظر رسید که تیواری به زبان «هریانوی» (Harianvi) روان. یک گویش هندی. در حال صحبت کردن است. «نیلکانت باکشی» (Neelkant Bakshi)، مدیر کمپین حزب بهاراتیا جاناتا به رسانه «وایس» (Vice) گفت: «فناوری جعل عمیق به ما کمک کرده که تلاش های کمپین انتخاباتی را به نحوی بی سابقه گسترش دهیم.» به همان روشی که سینتسیا با استفاده از هوش مصنوعی خود توانست هنگام تبلیغ یک کمپین آموزشی برای یک کار خیریه جهت ریشه کنی بیماری مالاریا، «دیوید بکام» فوتبالیست را در حال صحبت به ۹ زبان مختلف نشان دهد.

ریپاربلی می گوید: «بیشتر کار سینتسیا در آینده نزدیک روی تبدیل متن به ویدئوهای اطلاعاتی متمرکز خواهد بود.» وی توضیح می دهد: «تصور کنید شما مدیر شرکت بزرگ خرده فروشی با تعداد زیادی نیرو هستید و یک سند پنج صفحه ای پی دی اف در مورد دستورالعمل های امنیتی جدید ارسال می کنید. این شیوه، روش مؤثری برای برقراری ارتباط با آنها نیست.» چراکه هوش مصنوعی می تواند این اسناد را به یک فیلم جذاب تبدیل کند. «ما خودمان را جایگزین مناسبی برای ویدئوهای مرسوم نمی دانیم؛ این، چیزی نیست که شما برای صرفه جویی در هزینه تولید یک فیلم خریداری کنید. یک

جایگزین برای متن است.»

ریپاربلی سریعاً به این نکته اشاره می‌کند که شبکه‌های ساختگی نمی‌توانند احتمالاً در همه موقعیت‌ها فعالیت کنند. به عنوان نمونه، بعید است کسی به یک جلسه رودررو با رئیس خود که در واقع شامل گفت‌وگو با یک ویدئوی ساختگی ایجاد شده توسط یک چت‌بات است، علاقه‌مند شود. وی همچنین فکر می‌کند شفافیت ضروری است: «فکر می‌کنم در بسیاری از مکان‌هایی که باید ویدئوهای ساختگی به شما ارائه شود، مشخص خواهد بود که ویدئو ساختگی است.» اگر به طور کلی بتوان به این شفافیت دست یافت، برخی جنبه‌های در حال حاضر مشکل ساز، از جانب شبکه‌های ساختگی، به طور قابل توجهی کاهش خواهند یافت.

بهره‌اندیشه‌ها

برای سفارش اینترنتی این کتاب به وبسایت
فروشگاه انتشارات راه پرداخت
way2pay.shop



انتشارات راه پرداخت

مهم‌ترین فناوری همه‌منظوره عصر فعلی، هوش مصنوعی است. با اینکه اهمیت هوش مصنوعی امروزه بر همه آشکار است، اما بسیاری از افراد معنای آن را به اشتباه متوجه شده‌اند. مدیران، هوش مصنوعی را نوعی فناوری کلیدی ساختار شکن می‌دانند، کارمندان به عنوان یک ویران‌کننده شغل به آن می‌نگرند، مشاوران همیشه در صحبت‌هایشان به عنوان راه‌حلی همه‌جانبه از آن یاد می‌کنند و رسانه‌ها مرتب آن را بزرگ و ترسناک جلوه می‌دهند.

هوش مصنوعی هرچند طی دهه اخیر راه خود را در میان صنایع مختلف باز کرده است اما همچنان نظریه پردازان معتقدند که این فناوری جای رشد بسیار بیشتری دارد و از پتانسیلی برخوردار است که باید سال‌ها بگذرد که ظرفیت آن شکوفا شود.

مت برجس در این کتاب همه اطلاعاتی را که برای آشنایی با هوش مصنوعی لازم دارید را در اختیاران می‌گذارد و شیوه عملکرد آن‌ها را به شما توضیح می‌دهد. او به ثمره استفاده از این فناوری در محصولات مانند نرم‌افزار تشخیص صدا و اتومبیل‌های خودران اشاره می‌کند و پتانسیل آن برای ایجاد تغییرات انقلابی در همه ابعاد زندگی ما را بررسی می‌کند. همچنین این کتاب به جنبه‌های تاریک فناوری یادگیری ماشین هم می‌پردازد. حساسیت در برابر هک شدن، تمایل به سوگیری نسبت به برخی گروه‌های خاص و سوءاستفاده از آن توسط دولت‌ها از جمله نقاط ضعف فناوری یادگیری ماشین است که در این کتاب مورد بررسی قرار گرفته است. در آخر هم به یک سوال اصلی پاسخ داده می‌شود: آیا ماشین‌ها می‌توانند به اندازه انسان‌ها هوشمند شوند؟



کتاب هوش مصنوعی
با حمایت شرکت
آدونیس منتشر شده
و ناشر اصلی آن
انتشارات وایرد است

WIRED

ISBN 978-622-7702-04-0



۶۰ هزار تومان

انتشارات راه‌پروا

ناشر فناوری و نوآوری

way2pay.press