

مزایا و خطرات هوش مصنوعی ۱۱ ژانویه ۲۰۱۷

نویسنده:

رواتی کومار

فهرست مطالب

هوش مصنوعی چیست؟

چرا باید ایمنی هوش مصنوعی را مطالعه کنیم؟

هوش مصنوعی چگونه می‌تواند خطرناک باشد؟

چرا اخیراً علاقه به ایمنی هوش مصنوعی تا این حد افزایش یافته است؟

رایج‌ترین افسانه‌ها در مورد هوش مصنوعی پیشرفته

فرصیات مربوط به جدول زمانی

افسانه‌هایی درباره مناظره

افسانه‌هایی درباره خطرات هوش مصنوعی فرا بشری

هوش مصنوعی چیست؟

هوش مصنوعی (AI) به سرعت در حال پیشرفت است، از SIRI گرفته تا ماشین‌های خودران. اگرچه اغلب در داستان‌های علمی تخیلی به عنوان ربات‌هایی با ویژگی‌های انسانی به تصویر کشیده می‌شوند، اما هوش مصنوعی می‌تواند طیف وسیعی از چیزها را در بر بگیرد، از الگوریتم‌های جستجوی گوگل گرفته تا واتسون IBM و سلاح‌های رباتیک خودکار.

هوش مصنوعی که امروزه می‌شناسیم، به طور دقیق‌تر، هوش مصنوعی محدود (یا هوش مصنوعی ضعیف) است - که برای انجام وظایف خاص طراحی شده است (مثلاً فقط تشخیص چهره، فقط جستجوهای اینترنتی، فقط رانندگی). با این حال، هدف بلندمدت بسیاری از محققان، ایجاد هوش مصنوعی به معنای وسیع آن (AGI یا هوش مصنوعی قوی) است. در حالی که هوش مصنوعی با تعریف محدود، در هر کار خاصی، مانند بازی شطرنج یا حل معادلات، از انسان‌ها بهتر عمل می‌کند، هوش مصنوعی عمومی تقریباً در تمام وظایف شناختی از انسان‌ها بهتر عمل خواهد کرد.

چرا باید ایمنی هوش مصنوعی را مطالعه کنیم؟

در کوتاه‌مدت، هدف حفظ تأثیر مفید هوش مصنوعی برای جامعه، تحقیقات را در بسیاری از زمینه‌ها، از اقتصاد و حقوق گرفته تا مباحث فنی مانند تأیید، اعتبار، امنیت و کنترل، هدایت خواهد کرد. اگر لپ‌تاپ شما خراب شود یا هک شود، چیزی بیش از یک ناراحتی جزئی است، اما اگر سیستم‌های هوش مصنوعی، خودروها، هواپیماها، دستگاه‌های تنظیم ضربان قلب، سیستم‌های معاملاتی خودکار یا شبکه‌های برق شما را کنترل کنند اوضاع فرق می‌کند. اهمیت کار کردن این سیستم‌ها به روشی که شما می‌خواهید، بیشتر هم خواهد شد. یکی دیگر از چالش‌های کوتاه‌مدت، جلوگیری از یک مسابقه تسلیحاتی ویرانگر با سلاح‌های رباتیک خودکار است.

در درازمدت، یک سوال کلیدی این است که اگر در ایجاد هوش مصنوعی قوی موفق شویم و سیستم‌های هوش مصنوعی در تمام وظایف شناختی از انسان‌ها پیشی بگیرند، چه اتفاقی خواهد افتاد. همانطور که آی. جی. گود در سال ۱۹۶۵ اشاره کرد که طراحی سیستم‌های هوش مصنوعی هوشمندتر، خود یک وظیفه شناختی است. چنین سیستم‌هایی می‌توانند به‌طور مداوم خود را بهبود بخشند و تحولات فکری انفجاری را فراتر از هوش انسانی تجربه کنند. با استفاده از چنین هوش باورنکردنی برای اختراع فناوری‌های جدید انقلابی، می‌توانیم جنگ، بیماری و فقر را ریشه‌کن کنیم و ایجاد هوش مصنوعی قوی را به بزرگترین رویداد تاریخ بشر تبدیل کنیم. اما برخی از کارشناسان ابراز نگرانی کرده‌اند که هوش مصنوعی می‌تواند پایان بشریت باشد، مگر اینکه یاد بگیریم قبل از اینکه هوش مصنوعی به یک ابرهوشمند تبدیل شود، اهداف خود را با آن همسو کنیم.

برخی این سوال را مطرح می‌کنند که آیا هوش مصنوعی قوی هرگز محقق خواهد شد یا خیر، در حالی که برخی دیگر استدلال می‌کنند که ایجاد یک هوش مصنوعی فوق هوشمند تضمین شده است که مفید خواهد بود. FLI هر دوی این احتمالات را تشخیص می‌دهد، اما همچنین پتانسیل سیستم‌های هوش مصنوعی برای ایجاد آسیب‌های بزرگ، چه عمدی و چه غیرعمدی، را نیز در نظر می‌گیرد. ما معتقدیم که تحقیقات امروز می‌تواند به ما در آماده شدن برای پیامدهای وخیم فردا و جلوگیری از آنها کمک کند و به ما در اجتناب از خطرات و بهره‌مندی از مزایای هوش مصنوعی یاری رساند.

هوش مصنوعی چگونه می‌تواند خطرناک باشد؟

اکثر محققان معتقدند که بعید است یک هوش مصنوعی فوق هوشمند احساسات انسانی مانند عشق یا نفرت را از خود نشان دهد و هیچ دلیلی وجود ندارد که انتظار داشته باشیم یک هوش مصنوعی عمداً خیرخواه یا بدخواه باشد. در عوض، متخصصان هنگام بررسی اینکه هوش مصنوعی چگونه می‌تواند خطرناک باشد، دو سناریو را محتمل‌ترین می‌دانند:

هوش مصنوعی برای انجام کاری ویرانگر برنامه‌ریزی شده است: سلاح‌های رباتیک خودمختار، سیستم‌های هوش مصنوعی هستند که برای کشتن برنامه‌ریزی شده‌اند. در دست افراد ناباب، این سلاح‌ها به راحتی می‌توانند تلفات جمعی ایجاد کنند. علاوه بر این، یک مسابقه تسلیحاتی هوش مصنوعی می‌تواند ناخواسته منجر به جنگ‌های هوش مصنوعی با تلفات گسترده شود. برای جلوگیری از خرابکاری دشمن، این سلاح‌ها طوری طراحی خواهند شد که «خاموش کردن» آنها به سادگی بسیار دشوار باشد. بنابراین احتمال زیادی وجود دارد که انسان‌ها کنترل خود را بر چنین موقعیت‌هایی از دست بدهند. این خطر به معنای محدود کلمه به هوش مصنوعی نیز مربوط می‌شود، اما با هوشمندتر و خودمختارتر شدن هوش مصنوعی، این خطر نیز افزایش می‌یابد.

هوش مصنوعی برای انجام کاری مفید برنامه‌ریزی شده است، اما روش‌های مخربی را برای رسیدن به آن هدف توسعه می‌دهد: این سناریویی است که اگر اهداف ما کاملاً با اهداف هوش مصنوعی همسو نباشند، می‌تواند اتفاق بیفتد و جلوگیری از آن بسیار دشوار است. اگر از یک ماشین خودران مطیع و هوشمند بخواهید که در سریع‌ترین زمان ممکن به فرودگاه برسد، ممکن است درخواست شما را به معنای واقعی کلمه تشخیص دهد و با تمام سرعت به سمت فرودگاه حرکت کند، حتی اگر برایش مهم نباشد که توسط یک هلیکوپتر تعقیب شود یا با استفراغ پوشیده شود، حتی اگر این چیزی نباشد که واقعاً می‌خواستید. اگر به یک سیستم فوق هوشمند، یک پروژه مهندسی زمین‌شناسی بلندپروازانه محول شود، ممکن است به عنوان یک عارضه جانبی، اکوسیستم‌ها را ویران کند و تلاش‌های انسان برای متوقف کردن آن را به عنوان یک تهدید ببیند.

همانطور که این مثال‌ها نشان می‌دهند، نگرانی‌ها در مورد هوش مصنوعی پیشرفته، از سر توانایی است، نه از روی بدخواهی. هوش مصنوعی فوق هوشمند توانایی بالایی در دستیابی به اهداف خود خواهد داشت، اما اگر این اهداف با اهداف ما همسو نباشند، برای ما مشکل‌ساز خواهد شد. احتمالاً شما از آن دسته آدم‌های بدجنسی نیستید که از روی بدخواهی مورچه‌ها را له می‌کنند، اما اگر مسئول یک پروژه انرژی سبز برق‌آبی بودید و در منطقه‌ای که قرار بود سیل به آن وارد شود، لانه مورچه‌ای وجود داشت، مورچه‌ها را حذف می‌کردید. **هدف اصلی تحقیقات ایمنی هوش مصنوعی این است که هرگز انسان‌ها را در موقعیت مورچه‌ها قرار ندهیم.**

چرا اخیراً علاقه به ایمنی هوش مصنوعی تا این حد افزایش یافته است؟

بسیاری از نام‌های بزرگ در علم و فناوری، از جمله استیون هاوکینگ، ایلان ماسک، استیو وزنیاک و بیل گیتس، اخیراً نگرانی‌های خود را در مورد خطرات ناشی از هوش مصنوعی در رسانه‌ها یا از طریق نامه‌های سرگشاده ابراز کرده‌اند و بسیاری از محققان برجسته هوش مصنوعی نیز با این موضوع موافق بوده‌اند. چرا این موضوع ناگهان اینقدر مورد توجه قرار گرفته است؟

این ایده که تلاش برای دستیابی به هوش مصنوعی قوی سرانجام به موفقیت خواهد رسید، قرن‌ها موضوع داستان‌های علمی تخیلی محسوب می‌شد. اما به لطف پیشرفت‌های اخیر، بسیاری از نقاط عطف هوش مصنوعی که کارشناسان تنها پنج سال پیش تصور می‌کردند دهه‌ها با ما فاصله دارند، اکنون محقق شده‌اند و بسیاری از کارشناسان اکنون به طور جدی احتمال ظهور ابرهوش را در طول عمر نسل ما در نظر می‌گیرند. در حالی که برخی از کارشناسان هنوز حدس می‌زنند که هوش مصنوعی در سطح انسان قرن‌ها با ما فاصله دارد، اکثر متخصصان هوش مصنوعی در کنفرانس پورتوریکو در سال ۲۰۱۵ تخمین زدند که این اتفاق قبل از سال ۲۰۶۰ رخ خواهد داد. با توجه به اینکه ممکن است دهه‌ها طول بکشد تا مطالعات ایمنی لازم انجام شود، عاقلانه است که تحقیقات را از همین حالا شروع کنیم.

از آنجا که هوش مصنوعی پتانسیل پیشی گرفتن از هوش هر انسانی را دارد، ما هیچ راه مطمئنی برای پیش‌بینی رفتار آن نداریم. ما نمی‌توانیم از پیشرفت‌های تکنولوژیکی گذشته به عنوان پایه و اساس استفاده کنیم، زیرا هرگز چیزی با توانایی فریب دادن ما، چه عمداً و چه ناخودآگاه، نساخته‌ایم. شاید بهترین مثال از اینکه با چه امکاناتی روبرو هستیم، تکامل خودمان باشد. **بشریت در حال حاضر بر زمین تسلط دارد، نه به این دلیل که ما قوی‌ترین، سریع‌ترین یا بزرگترین هستیم، بلکه به این دلیل که ما باهوش‌ترین هستیم. اگر ما دیگر باهوش‌ترین نباشیم، این تسلط ممکن است دوام نیاورد.**

موضع FLI این است که تمدن ما تا زمانی شکوفا خواهد شد که بتوانیم از نیروهای رو به رشد فناوری پیشی بگیریم و بر آن غلبه کنیم. موضع FLI در مورد فناوری هوش مصنوعی این است که بهترین راه برای پیروزی در این رقابت، مانع تراشی در مورد اول نیست، بلکه تسریع در مورد دوم با حمایت از تحقیقات در زمینه ایمنی هوش مصنوعی است.

رایج‌ترین افسانه‌ها در مورد هوش مصنوعی پیشرفته

گفتگوهای جذابی در مورد آینده هوش مصنوعی و معنای آن، یا آنچه که باید برای بشریت داشته باشد، در حال انجام است. برخی از این بحث‌های جذاب حتی متخصصان برجسته جهان را نیز دچار اختلاف نظر می‌کند. برای مثال، تأثیر آینده هوش مصنوعی بر بازار کار، اینکه آیا هوش مصنوعی در سطح انسان توسعه خواهد یافت و اگر چنین شود، چه زمانی این اتفاق خواهد افتاد، اینکه آیا منجر به انفجار اطلاعات خواهد شد و اینکه آیا باید از آن استقبال کنیم یا از آن بترسیم. اما شبیه‌مناظره‌های خسته‌کننده‌ای هم وجود دارند که از سوء تفاهم‌ها ناشی می‌شوند. بیایید برخی از رایج‌ترین فرضیات را روشن کنیم تا بتوانیم به جای تصورات غلط، بر بحث‌های جالب و سوالات باز تمرکز کنیم.

فرضیات مربوط به جدول زمانی

اول، چقدر طول خواهد کشید تا ماشین‌ها به طور قابل توجهی از سطح هوش انسان پیشی بگیرند؟ در اینجا چند فرضیه در مورد جدول زمانی آورده شده است: یک تصور غلط رایج این است که ما می‌توانیم به این سوال پاسخ دقیقی بدهیم.

یک فرض رایج این است که ما مطمئناً در این قرن به هوش مصنوعی فوق بشری دست خواهیم یافت. در واقع، فناوری‌ها در گذشته اغلب بیش از حد مورد توجه قرار گرفته‌اند. نیروگاه‌های همجوشی هسته‌ای و ماشین‌های پرنده که گمان می‌رفت تاکنون توسعه یافته‌اند، هنوز هیچ نشانه‌ای از آنها دیده نمی‌شود. هوش مصنوعی حتی توسط بنیانگذاران این حوزه نیز بیش از حد مورد توجه قرار گرفته است. برای مثال، جان مک‌کارتی (که اصطلاح «هوش مصنوعی» را ابداع کرد)، ماروین مینسکی، ناتانیل روچستر و کلود شانون پیش‌بینی بیش از حد خوش‌بینانه‌ی زیر را در مورد آنچه یک کامپیوتر ابتدایی می‌تواند در عرض دو ماه به دست آورد، نوشتند: «ما پیشنهاد می‌کنیم که در تابستان ۱۹۵۶ به مدت دو ماه در کالج دارتموث در مورد هوش مصنوعی تحقیق کنیم و ده نفر در آن مشارکت داشته باشند... ما سعی خواهیم کرد راه‌هایی برای ماشین‌ها پیدا کنیم تا از زبان استفاده کنند، انتزاعات و مفاهیم را شکل دهند، مشکلاتی را که در حال حاضر انسان‌ها با آن مواجه هستند حل کنند و راه‌هایی برای بهبود خود پیدا کنند. ما معتقدیم که یک گروه منتخب از دانشمندان که برای یک تابستان با هم کار می‌کنند، می‌توانند در یک یا چند مورد از این مشکلات پیشرفت قابل توجهی داشته باشند.»

در همین حال، یک تصور رایج مخالف این است که ما مطمئن هستیم که در این قرن هوش مصنوعی فوق بشری توسعه نخواهیم داد. محققان طیف گسترده‌ای از پیش‌بینی‌ها را در مورد زمان توسعه هوش مصنوعی فوق بشری ارائه می‌دهند، اما با توجه به سابقه ناامیدکننده پیش‌بینی‌های بدبینانه به فناوری تاکنون، نمی‌توانیم مطمئن باشیم که این اتفاق در این قرن رخ خواهد داد. برای مثال، در سال ۱۹۳۳، ارنست رادرفورد، که مسلماً بزرگترین فیزیکدان هسته‌ای زمان خود بود، توسعه انرژی اتمی را «شوخ‌بازی» خواند. با این حال، کمتر از ۲۴ ساعت بعد، زیلارد واکنش زنجیره‌ای هسته‌ای را اختراع کرد. و ریچارد وولی، ستاره‌شناس سلطنتی، در سال ۱۹۵۶ سفر بین سیاره‌ای را «مزخرف محض» خواند. افراطی‌ترین فرض این است که هوش مصنوعی فوق بشری از نظر فیزیکی غیرممکن است و هرگز توسعه نخواهد یافت. اما فیزیکدانان می‌دانند که مغز از کوارک‌ها و الکترون‌هایی تشکیل شده است که طوری کنار هم قرار گرفته‌اند که مانند یک کامپیوتر قدرتمند عمل کنند، و هیچ قانون فیزیکی وجود ندارد که ما را از ساختن توده‌ای از کوارک‌ها که از مغز هوشمندتر هستند، باز دارد.

چند سال دیگر حداقل ۵۰٪ احتمال دارد که هوش مصنوعی در سطح انسان را ببینیم؟

نظرسنجی‌های زیادی انجام شده و از محققان هوش مصنوعی پرسیده شده است که آیا فکر می‌کنند این فناوری توسعه خواهد یافت یا خیر. این مطالعات همگی به یک نتیجه رسیدند. متخصصان برجسته جهان با این نظر مخالفند، بنابراین ما هیچ ایده‌ای نداریم. برای مثال، در یک نظرسنجی از محققان هوش مصنوعی در کنفرانس هوش مصنوعی پورتوریکو در سال ۲۰۱۵، میانگین (میان) پاسخ ۲۰۴۵ بود، اما برخی از محققان پیش‌بینی کردند که این زمان صدها سال یا بیشتر در آینده خواهد بود.

یک فرض مرتبط این است که افرادی که نگران هوش مصنوعی هستند معتقدند که تنها چند سال با آن فاصله داریم. در واقع، اکثر افرادی که نگرانی خود را در مورد هوش مصنوعی فوق بشری اعلام کرده‌اند، فکر نمی‌کنند که حداقل تا چند دهه آینده توسعه پیدا کند. اما آنها استدلال می‌کنند که مگر اینکه بتوانیم ۱۰۰٪ مطمئن باشیم که هوش مصنوعی در این قرن اتفاق نخواهد افتاد، عاقلانه است که تحقیق در مورد اقدامات ایمنی را برای آمادگی در برابر این احتمال آغاز کنیم. بسیاری از مسائل ایمنی مرتبط با هوش مصنوعی در سطح انسان بسیار دشوار هستند و حل آنها می‌تواند دهه‌ها طول بکشد. بنابراین عاقلانه است که تحقیقات خود را از همین حالا شروع کنید، نه اینکه یک شبه وقتی چند برنامه‌نویس هنگام نوشیدن رد بول تصمیم می‌گیرند تغییر رویه بدهند، شروع کنید.

افسانه‌هایی درباره مناظره

یکی دیگر از تصورات غلط رایج این است که فقط کسانی که نگران هوش مصنوعی هستند و از تحقیقات ایمنی هوش مصنوعی حمایت می‌کنند، ضد نوآوری هستند و اطلاعات زیادی در مورد هوش مصنوعی ندارند. وقتی استوارت راسل، نویسنده کتاب درسی استاندارد هوش مصنوعی، این را در طول ارائه‌ای در پورتوریکو گفت، حضار از خنده منفجر شدند. یک تصور غلط مرتبط این است که حمایت از تحقیقات ایمنی هوش مصنوعی موضوعی بسیار بحث‌برانگیز است. در واقع، برای حمایت از سرمایه‌گذاری‌های اندک در تحقیقات ایمنی هوش مصنوعی، نیازی نیست که متقاعد شد که خطرات بالا هستند و نمی‌توان آنها را نادیده گرفت. همانطور که سرمایه‌گذاری اندک در بیمه خانه با احتمال غیرقابل اغماض آتش گرفتن خانه شما توجیه می‌شود.

یکی از عوامل این است که رسانه‌ها بحث ایمنی هوش مصنوعی را جنجالی‌تر از آنچه واقعاً هست جلوه می‌دهند. گذشته از همه اینها، ایجاد ترس سودآور است. مقالاتی که از نقل قول‌های خارج از متن برای اعلام نابودی قریب‌الوقوع استفاده می‌کنند، نسبت به مقالات ظریف و متعادل ترجیح داده می‌شوند. در نتیجه، دو نفر که فقط از طریق نقل قول‌های رسانه‌ها از مواضع یکدیگر مطلع هستند، ممکن است فکر کنند که بیش از آنچه در واقع هست، با هم اختلاف نظر دارند. برای مثال، یک فرد شکاک به فناوری که صرفاً درباره موضع بیل گیتس در یک روزنامه بریتانیایی می‌خواند، ممکن است فرض کند که گیتس به قریب‌الوقوع بودن هوش فرابشری اعتقاد دارد. به همین ترتیب، یک طرفدار هوش مصنوعی که چیزی در مورد موضع اندرو نگ، فراتر از نقل قولش در مورد افزایش جمعیت در مریخ، نمی‌داند، ممکن است فرض کند که او به ایمنی هوش مصنوعی اهمیتی نمی‌دهد، اما در واقع اشتباه می‌کند. نکته اصلی این است که تخمین‌های بلندمدت نگ به این معنی است که او طبیعتاً تمایل دارد مسائل هوش مصنوعی کوتاه‌مدت را نسبت به مسائل بلندمدت در اولویت قرار دهد.

افسانه‌هایی درباره خطرات هوش مصنوعی فرابشری

بسیاری از محققان هوش مصنوعی از تیتلر «استیون هاوکینگ هشدار می‌دهد که ظهور ربات‌ها می‌تواند برای بشریت فاجعه‌بار باشد» شگفت‌زده شده‌اند. و من مقالات مشابه بی‌شماری دیده‌ام. معمولاً این مقالات با تصویری از یک ربات بدجنس با سلاح همراه هستند و این‌طور القا می‌کنند که ما نگران این هستیم که ربات‌ها برخیزند، دارای شعور و/یا شرور شوند و ما را بکشند. به طور خلاصه، مقالاتی مانند این، سناریوهایی را که محققان هوش مصنوعی نیازی به نگرانی در مورد آنها ندارند، به طور مختصر و مفید خلاصه می‌کنند. این سناریو شامل سه تصور غلط جداگانه است: نگرانی در مورد آگاهی، نگرانی در مورد شر و نگرانی در مورد ربات‌ها.

انسان‌ها هنگام رانندگی، تجربیات ذهنی از صداها و رنگ‌ها دارند. اما آیا یک ماشین خودران یک تجربه ذهنی خواهد داشت؟ آیا یک ماشین خودران چیزی را حس خواهد کرد؟ این اسرار آگاهی، اگرچه به خودی خود جذاب هستند، اما به خطرات هوش مصنوعی ربطی ندارند. اگر با یک ماشین بدون راننده تصادف کنید، فرقی نمی‌کند که ماشین از نظر ذهنی هوشیار به نظر برسد یا نه. به همین ترتیب، کاری که یک هوش مصنوعی فوق هوشمند انجام می‌دهد بر ما به عنوان انسان تأثیر می‌گذارد، نه آنچه که ما به صورت ذهنی احساس می‌کنیم.

ترس از اینکه ماشین‌ها به موجوداتی شرور تبدیل شوند، یکی دیگر از عوامل حواس‌پرتی است. چیزی که واقعاً باید نگران‌ش باشیم، شایستگی است، نه بدخواهی. یک هوش مصنوعی فوق هوشمند مسلماً در دستیابی به اهداف خود، هر چه که باشند، بسیار خوب عمل خواهد کرد، بنابراین باید مطمئن شویم که اهداف آن با اهداف ما همسو است.

تصور غلط در مورد هوشیاری مربوط به این باور است که ماشین‌ها نمی‌توانند هدف داشته باشند. یک ماشین می‌تواند اهدافی داشته باشد، به معنای محدود کلمه، یعنی نشان دادن رفتاری کاملاً هدفمند. این موضوع رفتار موشک‌های جستجوگر حرارت را توضیح می‌دهد که به آنها اجازه می‌دهد تا با بیشترین کارایی به اهداف حمله کنند. اگر از جانب ماشینی که اهدافش با اهداف شما همسو نیست، احساس خطر می‌کنید، مشکل این نیست که آیا آن ماشین آگاه یا هدفمند است، بلکه مشکل اهدافی است که دارد، به معنای دقیق کلمه. اگر توسط یک موشک ردياب حرارتی تعقیب می‌شدید، احتمالاً فریاد نمی‌زدید: «ماشین‌ها نمی‌توانند هدف‌گیری کنند، بنابراین من چیزی برای نگرانی ندارم!»

با توجه به اینکه روزنامه‌نگاران بی‌وقفه روی ربات‌ها تمرکز می‌کنند و بسیاری از مقالات خود را با عکس‌هایی از هیولاهای فلزی بدقیافه با چشمان قرمز درخشان تزئین می‌کنند، من با رادنی بروکس و دیگر پیشگامان رباتیک که توسط مطبوعات زرد اهریمنی جلوه داده شده‌اند، همدردی می‌کنم. در واقع، نگرانی اصلی کسانی که هوش مصنوعی مفید را ترویج می‌دهند، ربات‌ها نیستند، بلکه خودِ هوش است. به طور خاص، اهداف هوش با اهداف ما همسو نیستند. چنین هوش فراانسانی ناسازگار برای ایجاد مشکل برای ما نیازی به یک بدن رباتی ندارد. تنها چیزی که نیاز دارید یک اتصال اینترنتی است. این هوش به تنهایی می‌تواند بازارهای مالی را تسخیر کند، از محققان انسانی پیشی بگیرد، در مانور دادن از رهبران انسانی پیشی بگیرد و سلاح‌هایی فراتر از درک ما بسازد. حتی اگر ساخت یک ربات از نظر فیزیکی غیرممکن بود، یک هوش مصنوعی فوق‌العاده هوشمند و ثروتمند می‌توانست به راحتی به تعداد زیادی از انسان‌ها پول بدهد یا آنها را دستکاری کند تا ناخودآگاه دستوراتش را انجام دهند.

یک تصور غلط رایج در مورد ربات‌ها این است که ماشین‌ها نمی‌توانند انسان‌ها را کنترل کنند. هوش، امکان کنترل را فراهم می‌کند. انسان‌ها ببرها را کنترل می‌کنند نه به این دلیل که ما قوی‌تر هستیم، بلکه به این دلیل که ما باهوش‌تر هستیم. این یعنی اگر ما دیگر باهوش‌ترین افراد روی کره زمین نباشیم، ممکن است دیگر اوضاع را تحت کنترل نداشته باشیم.

بحث جالبی بود؟

اگر وقت خود را صرف تصورات غلط بالا نمی‌کردیم، می‌توانستیم روی بحث‌های واقعاً جالبی که متخصصان را از هم جدا می‌کند، تمرکز کنیم. به چه آینده‌ای می‌خواهید برسید؟ آیا باید سلاح‌های خودمختار مرگبار توسعه داده شوند؟ به نظر شما اتوماسیون کسب و کار چگونه باید توسعه یابد؟ چه توصیه‌های شغلی به فرزندان‌تان می‌دهید؟ آیا ترجیح می‌دهید مشاغل جدید جایگزین مشاغل قدیمی شوند، یا جامعه‌ای که در آن همه از اوقات فراغت و ثروت تولید شده توسط ماشین بدون شغل لذت می‌برند؟ آیا می‌خواهید در آینده حیات فوق هوشمند ایجاد کنید و آن را در کیهان گسترش دهید؟ آیا ما ماشین‌های هوشمند را کنترل خواهیم کرد یا آنها ما را؟ آیا ماشین‌های هوشمند جایگزین ما خواهند شد، با ما همزیستی خواهند کرد یا با ما ادغام خواهند شد؟ انسان بودن در عصر هوش مصنوعی به چه معناست؟ می‌خواهید این به چه معنا باشد، و چگونه می‌توانیم آن آینده را به واقعیت تبدیل کنیم؟ به این بحث بپیوندید!

این محتوا اولین بار در تاریخ ۱۱ ژانویه ۲۰۱۷ در futureoflife.org منتشر شد.

درباره موسسه زندگی آینده

موسسه زندگی آینده (FLI) یک اندیشکده جهانی با تیمی متشکل از بیش از ۲۰ کارمند تمام وقت است که در سراسر ایالات متحده و اروپا فعالیت می‌کند. FLI از زمان تأسیس خود در سال ۲۰۱۴، در تلاش بوده است تا توسعه فناوری‌های دگرگون‌کننده را به سمت سودمندی برای زندگی و دوری از خطرات بسیار بزرگ سوق دهد. درباره مأموریت ما بیشتر بدانید یا کار ما را بررسی کنید.