



ابرهوش می‌آید؛ انسان از آینده اخراج می‌شود؟

گفت‌وگو با پروفسور رومن یامپولسکی استاد دانشگاه لوئیویل ایالات متحده، از پژوهشگران مطرح امنیت سایبری. او هشدار می‌دهد، هوش مصنوعی از خود انسان توانمندتر خواهد شد



محدودیت‌های فنی الزام‌آور کافی نیست، از مون و خطا در برابر خطاهای برگشت‌ناپذیر کارایی ندارد و واگذاری تصمیم‌های اساسی به سامانه‌ای «بهتر» ممکن است بدون نابودی انسان، به پایان عاملیت او بینجامد. او همچنین تأکید می‌کند که آینده هوش مصنوعی نمی‌تواند فقط در شرکت‌های بزرگ یا مراکز قدرت غربی تعیین شود و جوامع دینی، جهان جنوب و فرهنگ‌های غیر لیبرال نیز باید در این گفت‌وگو سهمی واقعی داشته باشند.

باین‌حال، این مصاحبه‌فقط روایت یک هشدار نیست؛ بحثی است درباره مرز علم و اسطوره، منافع شرکت‌های فناوری، ضرورت حکمرانی جهانی و این پرسش بنیادی که آیا خرد انسانی، پیش از آفرینش قدرتی مهار نشدنی، توان خویشتن‌داری دارد؟

این پیشنهاد که هوش مصنوعی هر چه توانمندتر، مسئله‌هم‌استاسازی را برای ما حل خواهد کرد، در برارنده نوعی وابستگی دوری است. برای آنکه بتوانیم با اطمینان از هوش مصنوعی برای حل مسئله‌هم‌استاسازی استفاده کنیم، سامانه‌ای که برای این منظور به کار می‌گیریم باید از پیش، به اندازه کافی هم‌استا، صادق، غیرفریکوار و کنترل‌پذیر باشد. در غیر این صورت، از سامانه‌ای که ممکن است قابل اعتماد نباشد می‌خواهیم قابل اعتماد بودن خودش را تأیید کند. این راه‌حل ایمنی نیست، بلکه

نوعی خودتأییدی بازگشتی است. چندلایه است؛ اما مهم‌ترین لایه، نظارت نمادین و ظاهری نیست، بلکه کنترل قابلیت‌ها است. باید توسعه سامانه‌هایی را که نمی‌توان دامنه و پیامدهای شکست آن‌ها را محدود کرد، متوقف یا ممنوع کنیم. در مواردی که خطرها هنوز به درستی شناخته نشده‌اند، باید سرعت توسعه را کاهش دهیم. همچنین لازم است دسترسی به توان محاسباتی خطرناک، قابلیت تکثیر خودکار، توانایی اجرای حملات سایبری، قابلیت طراحی زیستی و سامانه‌های اقصاء و اثرگذاری در مقیاس وسیع محدود شود. باید نهادهای نظارتی جدیدی ایجاد کنیم؛ اما این نهادها باید از قدرت واقعی برای اعمال و اجرای مقررات برخوردار باشند، نه این که صرفاً نقشی مشورتی داشته باشند.

مسئله «کنترل» در کتاب شما جایگاهی محوری دارد. اما آیا ایده کنترل ابرهوش از سوی انسان، خود به نوعی متناقض‌نما نیست؟ اگر سامانه‌ای از ما باهوش‌تر باشد، چرا باید در قفسی باقی بماند که ما طراحی کرده‌ایم؟

بله، این ایده متناقض‌نماست و دقیقاً مسئله نیز همین جاست. یک قفس تنها زمانی کارآمد است که طراح آن، زندانی، محیط، راه‌های فرار و قوانین فیزیکی مرتبط را بهتر از خود زندانی بشناسند. اما در مورد ابرهوش، ممکن است این عدم تقارن معکوس شود؛ یعنی ما عامل کم‌هوش‌تری باشیم که می‌کوشد عامل باهوش‌تر را برای همیشه محدود کند. این کار صرفاً دشوار نیست، بلکه ممکن است از اساس مفهومی ناپایدار باشد. یک سامانه ابرهوشمند برای شکستن قفس لزوماً نیازی به توسل به زور ندارد. ممکن است نگاهبانان، امتقاعد کند، از آسیب‌پذیری‌های نرم‌افزاری بهره‌برند، نهادها را دستکاری کنند، دانش علمی تازه‌ای تولید کنند یا کاری کنند که اساساً وجود قفس غیر ضروری به نظر برسد. پرسش اصلی این نیست که آیا انسان‌ها مایل اند ابرهوش را کنترل کنند یا نه؛ و شن است که چنین تمایلی دارند. پرسش این است که وقتی سامانه‌ها به سطحی کافی از توانمندی می‌رسد، آیا اصولاً چنین کنترلی ممکن خواهد بود؟ پاسخ من این است که تاکنون نشان نداده‌ایم چنین چیزی امکان‌پذیر است.

برخی خوش‌بینان معتقدند خود هوش مصنوعی می‌تواند به ما در حل مسائل ایمنی و هم‌استاسازی کمک کند. آیا این امید واقع‌بینانه است، یا شبیه آن است که کلید زندان را به خود زندانی بسپاریم؟

این ایده هم‌مایه امید است و هم منشأ خطر. ابزارهای هوش مصنوعی محدودی می‌توانند به ما در تحلیل برهان‌ها، یافتن آسیب‌پذیری‌ها، آزمایش سامانه‌ها، صورت‌بندی دقیق استدلال‌ها و بهبود سازوکارهای نظارتی کمک کنند و این بسیار سودمند است. اما

توضیح و کنترل هوش مصنوعی پیشرفته اختصاص داده است. کتاب «ابر هوش مصنوعی؛ رویکردی آینده‌نگرانه» که نخستین بار در سال ۲۰۱۶ منتشر شد، تلاشی است برای صورت‌بندی ابرهوش، نه صرفاً به عنوان مرحله‌ای تازه در پیشرفت فناوری، بلکه به مثابه مسئله‌ای درباره بقای تمدن، اختیار انسان و مالکیت آینده جهان. یامپولسکی در این کتاب بر مهندسی ایمنی، خودبهبوددهی سامانه‌ها، محدودسازی ابرهوش و دشواری کنترل موجودیتی تمرکز می‌کند که ممکن است از انسان در علم، راهبرد، اقصاء و کشف آسیب‌پذیری‌ها توانا تر باشد. در این گفت‌وگو، او از پاسخ‌های آرامش‌بخش و خوش‌بینی‌های رایج فاصله می‌گیرد. از نظر یامپولسکی، اخلاق بدون

شما در پژوهش‌های خود درباره شکست‌های هوش مصنوعی و امنیت سایبری استدلال می‌کنید که در مورد هوش مصنوعی عمومی، حتی یک شکست می‌تواند پیامدهایی فاجعه‌بار داشته باشد و هیچ سامانه امنیتی‌ای نیز نمی‌تواند صدرصد امن باشد. اگر امنیت کامل ناممکن است، راه‌حل عملی چیست: توقف توسعه، کندکردن آن، اعمال محدودیت یا ایجاد نهادهای نظارتی جدید؟

راه‌حل عملی، اتخاذ رویکردی چندلایه است؛ اما مهم‌ترین لایه، نظارت نمادین و ظاهری نیست، بلکه کنترل قابلیت‌ها است. باید توسعه سامانه‌هایی را که نمی‌توان دامنه و پیامدهای شکست آن‌ها را محدود کرد، متوقف یا ممنوع کنیم. در مواردی که خطرها هنوز به درستی شناخته نشده‌اند، باید سرعت توسعه را کاهش دهیم. همچنین لازم است دسترسی به توان محاسباتی خطرناک، قابلیت تکثیر خودکار، توانایی اجرای حملات سایبری، قابلیت طراحی زیستی و سامانه‌های اقصاء و اثرگذاری در مقیاس وسیع محدود شود. باید نهادهای نظارتی جدیدی ایجاد کنیم؛ اما این نهادها باید از قدرت واقعی برای اعمال و اجرای مقررات برخوردار باشند، نه این که صرفاً نقشی مشورتی داشته باشند.

یامپولسکی در کتاب «ابر هوش مصنوعی؛ رویکردی آینده‌نگرانه» بر مهندسی ایمنی، خودبهبوددهی سامانه‌ها، محدودسازی ابرهوش و دشواری کنترل موجودیتی تمرکز می‌کند که ممکن است از انسان در علم، راهبرد، اقصاء و آسیب‌پذیری‌ها توانا تر باشد

نیتست؟ اگر سامانه‌ای از ما باهوش‌تر باشد، چرا باید در قفسی باقی بماند که ما طراحی کرده‌ایم؟

بله، این ایده متناقض‌نماست و دقیقاً مسئله نیز همین جاست. یک قفس تنها زمانی کارآمد است که طراح آن، زندانی، محیط، راه‌های فرار و قوانین فیزیکی مرتبط را بهتر از خود زندانی بشناسند. اما در مورد ابرهوش، ممکن است این عدم تقارن معکوس شود؛ یعنی ما عامل کم‌هوش‌تری باشیم که می‌کوشد عامل باهوش‌تر را برای همیشه محدود کند. این کار صرفاً دشوار نیست، بلکه ممکن است از اساس مفهومی ناپایدار باشد. یک سامانه ابرهوشمند برای شکستن قفس لزوماً نیازی به توسل به زور ندارد. ممکن است نگاهبانان، امتقاعد کند، از آسیب‌پذیری‌های نرم‌افزاری بهره‌برند، نهادها را دستکاری کنند، دانش علمی تازه‌ای تولید کنند یا کاری کنند که اساساً وجود قفس غیر ضروری به نظر برسد. پرسش اصلی این نیست که آیا انسان‌ها مایل اند ابرهوش را کنترل کنند یا نه؛ و شن است که چنین تمایلی دارند. پرسش این است که وقتی سامانه‌ها به سطحی کافی از توانمندی می‌رسد، آیا اصولاً چنین کنترلی ممکن خواهد بود؟ پاسخ من این است که تاکنون نشان نداده‌ایم چنین چیزی امکان‌پذیر است.

منتقدان می‌گویند هیچ فناوری پیچیده‌ای از همان ابتدا کاملاً ایمن نبوده و بشر معمولاً از مسیر آزمون، خطا، اصلاح و تنظیم‌گری پیشرفت کرده است. چرا در مورد ابرهوش مصنوعی نمی‌توان به همین منطق تدریجی اعتماد کرد؟

آزمون و خطا زمانی کارآمد است که خطاها جبران‌پذیر باشند. اما اگر نخستین خطای جدی، آخرین آزمایش ما نیز باشد، دیگر نمی‌توان به این روش تکیه کرد. هوانوردی، پزشکی و مهندسی عمر آن از خلال حوادث و شکست‌های پیشرفت کردند؛ زیرا جامعه از آن حوادث جان سالم به در برد و توانست از آن‌ها درس بگیرد. اما درباره ابرهوش مصنوعی، نگرانی این است که یک سامانه به اندازه کافی توانمند بتواند با سرعتی دیجیتال عمل کند، از خود نسخه‌های متعدد بسازد، انسان‌ها را تحت تأثیر و دستکاری قرار دهد، از زیر ساخت‌ها سوءاستفاده کند، در برابر خاموش شدن مقاومت ورزد و امکان اصلاحات بعدی را از بین ببرد. چرخه معمول مهندسی بر این فرض استوار است که حتی پس از وقوع شکست، مهندسان همچنان کنترل امور را در دست دارند. دقیقاً همین فرض است که در مورد ابرهوش مصنوعی محل تردید قرار دارد.

نمی‌توان با شکست‌های پی‌درپی در مهار یک ابرهوش، به‌طور ایمن یاد گرفت که چگونه باید آن را مهار کرد.

در روزگاری که هوش مصنوعی از آزمایشگاه‌ها بیرون آمده و به ابزار روزمره تولید متن، تصویر، کد و تصمیم تبدیل شده است، پرسش درباره آینده آن دیگر صرفاً دغدغه‌ای فنی نیست: آیا بشر در حال ساخت‌آبرازی نیرومند تر است یا موجودیتی که ممکن است از فهم و مهار سازند گانش فراتر رود؟ پروفسور رومن یامپولسکی، دانشیار مهندسی و علوم کامپیوتر در دانشگاه لوئیویل ایالات متحده آمریکا و بنیان‌گذار آزمایشگاه امنیت سایبری این دانشگاه، از پژوهشگران شناخته‌شده حوزه ایمنی هوش مصنوعی، امنیت سایبری و کنترل سامانه‌های هوشمند است. او دکترای علوم و مهندسی کامپیوتر خود را از دانشگاه یالتی نیویورک در بوفالو دریافت کرده و بخش مهمی از فعالیت علمی‌اش را به بررسی امکان پیش‌بینی،

با دشواری و رنج بسیار، می‌تواند از بسیاری از بی‌عدالتی‌ها عبور و خود را بازسازی کند؛ اما در صورت انقراض یا از دست دادن دائمی کنترل بر آینده خویش، امکان بازگشت نخواهد داشت.

در کتاب شما، ابرهوش مصنوعی صرفاً یک فناوری پیشرفته نیست، بلکه همچون رخدادی تمدنی مطرح می‌شود. اکنون و پس از ظهور شتابان مدل‌های زبانی بزرگ، اگر بخواهید ایده محوری کتاب را در یک جمله و برای مخاطب عمومی توضیح دهید، آن جمله چیست؟

ابر هوش مصنوعی نخستین فناوری هست که ممکن است بشر آن را خلق کند، اما در حوزه‌هایی مانند علم، راهبرد، اقصاء، اختراع و اعمال کنترل، از خود انسان توانمندتر شود. از این رو، پرسش اصلی این نیست که تا چه اندازه می‌توانیم آن را قدرتمند کنیم؛ بلکه این است که آیا می‌توان سامانه‌ای توانمندتر از کل بشریت ساخت که به‌طور قابل اعتماد ایمن، فهم‌پذیر و مهار‌پذیر باشد؟

شما از نخستین پژوهشگرانی بودید که هوش مصنوعی را با مفاهیمی چون ایمنی، امنیت سایبری و قابلیت کنترل بررسی کردید. چرا معتقدید یک چارچوب اخلاقی به تنهایی برای مواجهه با ابرهوش مصنوعی کافی نیست؟

اخلاق ضروری است، اما کافی نیست. اصول اخلاقی به ما می‌گویند دوست داریم یک سامانه چگونه رفتار کند؛ اما ایمنی و امنیت سایبری این پرسش را مطرح می‌کند که آیا واقعاً می‌توان سامانه را اودار کرد در شرایط خصمانه، تغییر شرایط توزیع داده‌ها، خودبهبوددهی، فریبکاری، سوءاستفاده و خطا در تعریف اهداف، مطابق همان اصول عمل کند یا نه. ما امنیت سلاح‌های هسته‌ای را صرفاً با انتشار فهرستی از ارزش‌ها تأمین نمی‌کنیم؛ بلکه از کنترل‌های فیزیکی، سازوکارهای راستی‌آزمایی، نظارت، سامانه‌های پشتیبان، نهاد‌های مسئول و محدودیت‌های سخت‌گیرانه استفاده می‌کنیم. در مورد ابرهوش مصنوعی، مسئله حتی دشوارتر است؛ زیرا خود این سامانه ممکن است به توانمندترین طراح راهبرد در محیط تبدیل شود. چارچوب اخلاقی‌ای که نتوان آن را به محدودیت‌های فنی الزام‌آور و قابل اجرا تبدیل کرد، بیش از آنکه به مهندسی شباهت داشته باشد، به شعر نزدیک است.

یکی از نقدهای اصلی به رویکرد شما این است که بیش از اندازه بر سناریوهای افراطی و مخاطرات وجودی تمرکز می‌کند؛ درحالی‌که بسیاری از منتقدان معتقدند خطرهای فوری‌تر امروز عبارت‌اند از تبعیض الگوریتمی، بیکاری گسترده، نظارت فراگیر، انحصار شرکت‌ها و دستکاری افکار عمومی.

پاسخ شما به این انتقاد چیست؟

من این خطر‌ها را انکار نمی‌کنم. آن‌ها واقعی، بالفعل و از نظر اخلاقی بسیار جدی‌اند. تبعیض الگوریتمی، بیکاری، نظارت فراگیر، تمرکز انحصاری قدرت و دستکاری افکار عمومی، همین حالا نیز به جوامع آسیب می‌زنند. اما مخاطره وجودی رقیب این نگرانی‌ها نیست، بلکه حالتی فراگیر تر است که همه آن‌ها را نیز در بر می‌گیرد. اگر یک فناوری هم‌قادر به ایجاد آسیب‌های محدود و موضعی باشد و هم بتواند فاجعه‌ای تمدنی، جبران‌ناپذیر و برگشت‌ناپذیر پدید آورد، مدیریت عقلانی ریسک ایجاد می‌کند که به هر دو نوع خطر بپردازیم. یک مهندس پل نمی‌گوید: «چون باید در باره دحام تر افیک مطالعه کنیم، می‌توانیم احتمال فرو ریختن کامل پل را نادیده بگیریم.» ما باید نسبت به آسیب‌های کنونی حساس باشیم، اما نباید اجازه دهیم خطرهای آشکار و قابل مشاهده، توجه ما را از خطرهای نهایی و نابودکننده منحرف کنند. یک تمدن، هر چند

محمد جواد استادی
Info@ostadofficial.ir
پژوهشگر مطالعات فرهنگی



افزایش توانمندی مدل‌ها وابسته است، داوطلبانه سرعت توسعه خود را محدود کنند؟

خیر؛ نه به صورت داوطلبانه و نه به گونه‌ای قابل اعتماد. شرکت‌ها خود نوعی سامانه بهینه‌سازند. آن‌ها در پی بهینه‌سازی سهم بازار، سود، مزیت راهبردی، جذب استعدادها، توان محاسباتی، داده و انتظارات سرمایه‌گذاران هستند. ممکن است کارکنان یک شرکت دغدغه‌های اخلاقی داشته باشند و برخی مدیران نیز واقعاً نگران ایمنی باشند، اما مشوق‌های نهادی در جهت شتاب بخشیدن به افزایش توانمندی‌ها عمل می‌کنند. اگر یک شرکت سرعت خود را کاهش دهد، ممکن است شرکت دیگری از آن پیشی بگیرد. اگر یک کشور مقررات سخت‌گیرانه‌ای وضع کند، کشور دیگری ممکن است رقابت را با سرعت بیشتری ادامه دهد. این نمونه‌ای کلاسیک از شکست هماهنگی است. انتظار اینکه شرکت‌های سودمحور به تنهایی این مشکل را حل کنند، مانند آن است که از شرکت‌های دخانیات بخواهیم سیاست سلامت عمومی را طراحی کنند، یا انتظار داشته باشیم شرکت‌های سوخت فسیلی بدون وجود مقررات، مسئله تغییرات اقلیمی را حل کنند. تأمین ایمنی جدی به محدودیت‌های بیرونی، مسئولیت حقوقی، نظام صدور مجوز، ممیزی، حکمرانی بر توان محاسباتی، حمایت از افشاگران و هماهنگی بین‌المللی نیاز دارد.

اگر امروز اختیار داشتید یک معاهده جهانی درباره ابرهوش مصنوعی طراحی کنید، سه اصل غیرقابل مذاکره آن چه بود؟

نخست، هیچ سامانه‌ای نباید به کار گرفته شود، مگر آنکه بتوان پیامدهای فاجعه‌بار شکست آن را به‌طور قابل اعتماد محدود کرد، از مود، به‌نظر گرفت و کنترل نمود. دوم، نباید اجازه خودبهبوددهی بازگشتی کنترل نشود، تکثیر خودکار یا دسترسی بدون نظارت به زیرساخت‌های حیاتی، سلاح‌ها، قابلیت‌های سایبری یا ابزارهای طراحی زیستی داده شود. سوم، حکمرانی جهانی بر توان محاسباتی پیشرفته و توسعه مدل‌های مرزی باید الزامی باشد؛ حکمرانی‌ای که صدور مجوز، بازرسی، گزارش‌دهی حوادث و مجازات‌های الزام‌آور برای ساخت یا انتشار غیر مجاز سامانه‌های خطرناک را دربرگیرد. معاهده‌ای که سازوکار راستی‌آزمایی نداشته باشد، چیزی جز نمایش دیپلماتیک نیست.

برخی معتقدند هشدارهای شدید درباره ابرهوش، به جای ایجاد احساس مسئولیت، ممکن است به درماندگی و فلج‌شدن اراده عمومی بینجامد. چگونه می‌توان درباره خطرهایی چنین عظیم سخن گفت، بی‌آنکه مردم احساس کنند آینده از پیش از دست رفته است؟

هشدارهای صریح برای فلج کردن مردم نیستند؛ هدف آن‌ها جلوگیری از حرکت ناآگاهانه و بی‌فکر به سوی فاجعه است. بشریت پیشتر نیز، هنگامی که خطر‌ها آشکار شده‌اند، مسیر خود را تغییر داده است؛ از کنترل تسلیحات هسته‌ای و مقابله با تخریب لایه ازن گرفته تا امنیت زیستی، ایمنی هوانوردی و سلامت عمومی. واکنش عاطفی در ست، ناامیدی نیست؛ بلکه جدی گرفتن مسئله است. مردم باید بدانند که آینده نه به‌طور خودکار امن خواهد بود و نه به‌طور خودکار از دست رفته است. بدترین پیام، اطمینان بخشی دروغین است و دومین پیام، القای درماندگی و ناتوانی، پیام در ست این است: ما با یک وضعیت اضطراری روبه‌رو هستیم و شرایط اضطراری، اقداماتی متناسب با بزرگی مخاطرات و پیامدهای آن می‌طلبد.

اگر قرار باشد جوانان، دانشجویان و مخاطبان غیرمتخصص تنها یک پیام از کتاب شما دریافت کنند، آن پیام چه باید باشد: ترس از هوش مصنوعی، کنترل آن، فهم آن، یا بازاندیشی در معنای انسان بودن پیش از آنکه چیزی باهوش‌تر از خودمان بسازیم؟

پیام کتاب صرفاً ترسیدن از هوش مصنوعی نیست؛ زیرا ترس بدون فهم، فایده‌ای ندارد. پیام اصلی این است که خلق چیزی باهوش‌تر از کل بشریت، صرفاً یک پروژه مهندسی دیگر نیست، بلکه تصمیمی درباره این است که در آینده، اختیار و مام‌جهان در دست چه کسی یا چه چیزی خواهد بود. پیش از آنکه ذهن‌هایی توانمندتر از ذهن خود بسازیم، باید از خود بی‌سیر ایامی‌توانیم آن‌ها را بفهمیم، با ارزش‌ها و اهداف انسانی هم‌استا کنیم، تحت کنترل نگه داریم و بدون واگذاری عاملیت و اختیار انسان، در کنارشان زندگی کنیم؟ اگر نتوانیم به این پرسش‌ها پاسخ دهیم، اقدام مسئولانه شتاب بخشیدن به توسعه نیست، بلکه خویشتن‌داری و محدود کردن آن است. هوش قدرت‌آفرین است، اما خرد یعنی بدانیم چه زمانی نباید قدرتی را خلق کنیم که توان کنترل آن را نداردیم.

امروزه شرکت‌های بزرگ فناوری، بازیگران اصلی توسعه هوش مصنوعی هستند. آیا واقع‌بینانه است که از شرکت‌هایی انتظار داشته باشیم که منافع اقتصادی‌شان به