



اندیشکده فرهنگ و توسعه

یادداشت



نویسنده: رابرت ای. مینینگ و فریال سعید

منبع انتشار: استیمسون

تاریخ انتشار: اگوست ۲۰۲۵

مسابقه هوش مصنوعی:
نویدها و خطرات آرمان‌شهرهای فناوری

مسابقه جهانی شتاب‌دهنده برای تسلط هوش مصنوعی (AI)، که ناشی از اعتقاد به اجتناب‌ناپذیری، نویدبخشی تجاری گسترده و فوریت ژئوپلیتیکی آن - به‌ویژه در رابطه با رقابت ایالات متحده با چین - است، نگرانی‌های جدی را ایجاد می‌کند. اگرچه پتانسیل تحول‌آفرین هوش مصنوعی غیرقابل انکار است، اما توسعه آن با توجه ناکافی به خطرات سیستمی، از جمله جابجایی گسترده نیروی کار، فرسایش عامل انسانی و حاکمیت ناکافی، پیش می‌رود. به این ترتیب، دیدگاه غالب «آرمان‌شهر فناوری» - که توسط فناوری‌گرایان و سرمایه‌گذاران خطرپذیر برجسته به عنوان ترکیبی از ایدئولوژی و یک مدل کسب‌وکار که خطرات نزولی را به نفع رشد نامحدود به حداقل می‌رساند، مورد حمایت قرار می‌گیرد - قابل نقد است.

نویسندگان با ترسیم قیاس‌های تاریخی با جنبش تکنوکراسی دهه ۱۹۳۰ و شرکت هند شرقی بریتانیا، بررسی می‌کنند که چگونه شرکت‌های بزرگ فناوری به طور فزاینده‌ای به عنوان یک بازیگر مستقل عمل می‌کنند و با پاسخگویی محدود، حکومت و جامعه را تغییر شکل می‌دهند. با وجود هشدارهای برجسته از سوی پیشگامان هوش مصنوعی و رهبران شرکت‌ها، هماهنگی جهانی همچنان پراکنده است و ایالات متحده در تنظیم مقررات و سرمایه‌گذاری ایمنی عقب مانده است. بدون اصلاح مسیر به سمت نوآوری انسان‌محور و حفاظت‌های جهانی قابل اجرا، هوش مصنوعی نه با اخلاق و دوراندیشی، بلکه با قدرت و سود شکل می‌گیرد.

اعتقاد به اجتناب‌ناپذیری هوش مصنوعی (AI)، وعده مزایای بی‌حد و حصر، و ترس از باخت در برابر چین یا یکدیگر، رقابت صنعت هوش مصنوعی آمریکا را به یک مسابقه دو مرحله‌ای و خشمگین برای تسلط بر هوش مصنوعی سوق می‌دهد. این مسابقه با توجه ناکافی به خطرات - که بسیاری از آنها توسط خود سازندگان هوش مصنوعی و رهبران شرکت‌ها مطرح شده‌اند - در حال شتاب گرفتن است. نگرانی‌ها در مورد ایمنی، تأثیر انسان و ناتوانی در جلوگیری از نتایج فاجعه‌بار در پی دستیابی به سرعت و برتری، کم‌اهمیت جلوه داده شده‌اند. چالش، متوقف کردن قطار یا حتی کند کردن آن نیست - بلکه تغییر جهت به گونه‌ای است که اطمینان حاصل شود نوآوری با اهداف انسانی همسو باقی می‌ماند.

آرمان‌شهر فناوری: «همه چیز بیشتر، برای همیشه»

سم آلتمن، مدیرعامل OpenAI، در یک پست وبلاگی در مورد هوش مصنوعی برتر، که به عنوان هوش مصنوعی فراتر از هوش انسانی تعریف می‌شود، توضیح می‌دهد: «ما در حال ساختن مغزی برای جهان هستیم.» ربات‌ها می‌توانند روزی با استخراج مواد خام، راندن کامیون‌ها و اداره کارخانه‌ها، کل زنجیره‌های تأمین را اداره کنند. حتی شگفت‌انگیزتر اینکه، برخی ربات‌ها می‌توانند ربات‌های بیشتری برای ساخت کارخانه‌های ساخت تراشه، مراکز داده و سایر زیرساخت‌ها تولید کنند. ماشین‌ها نه تنها آینده را تأمین می‌کنند، بلکه آن را خواهند ساخت - به طور نامحدود.

رهبران فناوری و سرمایه‌گذاران خطرپذیر مانند ایلان ماسک، مارک زاکربرگ و مارک اندریسن، آینده‌ای را تصور می‌کنند که توسط فناوری تعیین می‌شود، جایی که هوش مصنوعی بیماری را ریشه‌کن می‌کند، فروپاشی محیط زیست را معکوس می‌کند، انرژی بی‌حد و حصر را کشف می‌کند و عصری از فراوانی را آغاز می‌کند. چه چیزی را نباید دوست داشت؟ با این حال، رهبران فناوری همچنین تصور می‌کنند که هوش مصنوعی جامعه را چنان کاملاً متحول خواهد کرد که مردم به دنیای مجازی پناه می‌برند، دوستی‌های هوش مصنوعی تشکیل می‌دهند، با ارزهای دیجیتال معامله می‌کنند و به درآمد پایه جهانی که بخشی از آن توسط شرکت‌های فناوری تأمین می‌شود، متکی خواهند بود. برخی آن را به یک چشم‌انداز افراطی و تاریک‌تر می‌برند که در آن دولت-ملت‌ها به تدریج در دولت‌های شبکه‌ای که توسط نخبگان تکنوکرات اداره می‌شوند - نسخه‌ای مدرن از پادشاهان فیلسوف سقراط - حل می‌شوند. واقع‌بینانه یا نه، این [رویکرد] در جاه‌طلبی خود کاملاً دگرگون‌کننده است، بدون اینکه به اندازه کافی با پیامدهای منفی آن دست و پنجه نرم کند. این [رویکرد]، چشم‌انداز دیستوپیایی آلدوس هاکسلی در کتاب «دنیای قشنگ نو» را منعکس می‌کند و اعتقاد اصلی آن این است که فناوری، به ویژه هوش مصنوعی، «همه چیز بیشتر را برای همیشه» ارائه خواهد داد.

تکنوکراسی دوباره متولد شد





چنین جهانی یادآور جنبش تکنوکراسی دهه ۱۹۳۰ است که توسط پدر بزرگ ایلان ماسک، جاشوا ان. هالدمن، رهبری می‌شد و از حکومت مهندسان و دانشمندان حمایت می‌کرد، اما نتوانست توجه زیادی را به خود جلب کند. در نسخه جدید ۲۰۲۳ مارک اندریسن، با عنوان «مانیفست خوش‌بین به فناوری»، این سرمایه‌گذار خطرپذیر استدلال می‌کند که همه مشکلات، چه طبیعی و چه تکنولوژیکی، را می‌توان با فناوری بیشتر حل کرد. این طرز فکر در سال ۲۰۱۸ واکنش شدیدی را برانگیخت که ناشی از قدرت بی‌حد و حصر شرکت‌های فناوری بر داده‌های کاربران و محتوای پلتفرم‌هایشان بود و اعتماد مصرف‌کننده و این توهم را که فناوری بیشتر ناگزیر به پیشرفت منجر می‌شود، از بین برد.

اندریسن «دشمن» پیشرفت را دولت‌گرایی، جمع‌گرایی، برنامه‌ریزی مرکزی و انحصارگرایی می‌داند. این نیز آشناست: سرمایه‌داری بازار آزاد با حداقل مقررات، همان محیطی است که رشد اولیه اینترنت را تقویت کرد. مانیفست اندریسن با جامعه مانند یک برنامه کاربردی خراب رفتار می‌کند: ناقص، اما قابل اصلاح، تا زمانی که توسعه‌دهندگان به ارتقاء مجموعه فناوری ادامه دهند. به طور مشابه، دشمنانی که او شناسایی می‌کند در واقع نتایج پیچیده تاریخ بشر، مبارزات قدرت و ایدئولوژی‌های متضاد هستند. علاوه بر این، آنها نقص‌هایی نیستند که با یک وصله نرم‌افزاری بر آنها غلبه شود. آنها بخشی از شرایط انسانی هستند.

جایی که انسان‌ها در تصویر جای می‌گیرند، زمانی که ماشین‌ها تمام کارهای سنگین را انجام می‌دهند، به وضوح غایب است. طبق این دیدگاه، ماشین‌ها برای مردم فکر می‌کنند، می‌سازند و تصمیم می‌گیرند - نه با آنها. این آینده‌ای از اجماع مهندسی شده است، جایی که سیاست و رضایت حکومت‌شوندگان در کنار عامل انسانی ناپدید می‌شود. جای تعجب نیست که برخی در این صنعت، مانند ساندرا پیچای، مدیرعامل گوگل، از ادغام دانشمندان علوم اجتماعی، فیلسوفان و اخلاق‌گرایان در گفتگو حمایت می‌کنند - به ویژه برای یادآوری به دانشمندان داده و سرمایه‌گذاران خطرپذیر سیلیکون ولی که عامل انسانی باید دست نخورده باقی بماند. این بخشی از نرم‌افزار است که باید اصلاح شود.

آشفته‌گی اجتماعی و ظهور دولت فناوری

با این حال، اشتباه است که طرفداران آرمان‌شهرهای فناوری را صرفاً سودجویان یا ایدئولوگ‌ها بدانیم. آنها دلایل خوبی دارند که باور کنند هوش مصنوعی می‌تواند مشکلات قبلاً لاینحل را حل کند. آلفافولد (AlphaFold) شرکت دیپ‌ماینند (DeepMind) در حال حاضر زیست‌شناسی را متحول کرده و به خاطر پیش‌بینی دقیق ساختارهای پروتئین، جایزه نوبل ۲۰۲۴ را برای سازندگان‌شان به ارمغان آورده است، پیشرفتی که کشف دارو را تسریع می‌کند. هشدارهای زلزله گوگل، امنیت عمومی را بهبود می‌بخشد و خنک‌سازی مبتنی بر هوش مصنوعی می‌تواند مصرف انرژی تجاری را ۹ تا ۱۳ درصد کاهش دهد. هوش مصنوعی بدون شک پتانسیل عظیمی برای بهبود زندگی بشر دارد.

با این حال، هوش مصنوعی همچنین نوید می‌دهد که تمدن را با سرعتی بی‌سابقه متحول کند، در حالی که جوامع را مختل می‌کند. از نظر تاریخی، حتی در جوامعی با اوقات فراغت فراوان، مردم کار هدفمندی داشته‌اند. با این حال، صحبت از درآمد پایه جهانی خلاف این را نشان می‌دهد. هنگامی که انقلاب صنعتی جایگزین کار فیزیکی پیشاصنعتی و روستایی شد، مشاغل جدیدی ایجاد شد و تغییرات طی چند قرن آشکار شد و به جوامع فرصت داد تا خود را وفق دهند. در مقابل، انقلاب هوش مصنوعی، تهدید می‌کند که به سرعت جایگزین کارگران دانش‌محور شود، بدون اینکه هیچ شغل جایگزین فوری ارائه دهد. نیمی از مشاغل یقه سفید سطح پایه ممکن است ظرف پنج سال از بین بروند. فارغ التحصیلان امروزی دانشکده‌های بازرگانی نمی‌توانند بانکداران فردا شوند اگر مسیر کارآموزی به تخصص بدون چیزی برای جایگزینی آن از بین برود. خالی کردن طبقه حرفه‌ای با جایگزینی مشاغل با درآمد متوسط و دستمزد بالا با نقش‌های کمتر و با دستمزد پایین‌تر که صرفاً از سیستم‌های خودکار پشتیبانی می‌کنند، تحرک رو به بالا را از بین می‌برد. علاوه بر این، هنگامی که هوش مصنوعی مشاغلی را که به استدلال، قضاوت، پاسخگویی، اخلاق و وجدان متکی هستند - استعدادهای انسانی که فاقد آن هستند - خودکار می‌کند، عملکرد و ایمنی می‌تواند آسیب ببیند.

فشار بی‌وقفه برای سرعت، وظیفه اساسی آمادگی برای تغییر را تحت الشعاع قرار می‌دهد. هوش مصنوعی در حال حاضر در حال بازسازی جامعه است و نحوه ارتباط، کار، جنگ و مشارکت آمریکایی‌ها در سیاست را متحول می‌کند. این فناوری به لبه پیشرو در امنیتی‌سازی ملی تصمیم‌گیری‌های اقتصادی و رقابت ژئواستراتژیک تبدیل شده است. با جذب کاربران و گسترش سریع آن، دگرگونی جامعه سرعت می‌گیرد، به طوری که آنها باید خود را با آن وفق دهند. استقرار سریع هوش مصنوعی چین - در نظارت،





حکومتداری و کاربردهای نظامی - اغلب برای توجیه شتاب بیشتر مورد استناد قرار می‌گیرد. با این حال، اگر آمریکا در توسعه مدل‌های پیشرفته هوش مصنوعی پیشرو باشد در حالی که چین در استقرار مدل‌های کمتر پیشرفته خود پیشرو باشد، چالش فقط استقرار سریع‌تر نیست. چالش، رهبری با استقرار متفاوت است، یعنی با اولویت دادن به ایمنی، شفافیت و آمادگی برای اختلال، و در نتیجه تثبیت استقرار هوش مصنوعی آمریکا در اصول انسان‌محور به جای فوریت ژئوپلیتیکی.

ایان برمر، دانشمند علوم سیاسی آمریکایی، در مقاله‌ای که اخیراً در فارن افیرز منتشر کرده، این دنیای جدید را «تکنوقطبی» توصیف می‌کند، جایی که صنعت فناوری یک بازیگر غیردولتی و فراگیر است که ژئوپلیتیک و زندگی عمومی را شکل می‌دهد. بدون شک، شرکت‌های بزرگ فناوری ایالات متحده یا شرکت‌های «بزرگ فناوری»، نقش‌هایی را ایفا می‌کنند که از نظر تاریخی در سلطه دولت-ملت‌ها بوده‌اند و از برخی جهات، با دولت ادغام شده‌اند. اگرچه مشخص نیست که آیا آنها یک «قطب» جهانی مانند چین، روسیه یا اتحادیه اروپا تشکیل می‌دهند یا خیر، اما مسلم است که شرکت‌های بزرگ فناوری نمایانگر یک شکل جدید - عمدتاً غیرپاسخگو - از قدرت هستند.

شرکت‌های بزرگ فناوری، کمپانی هند شرقی بریتانیا را هدایت می‌کنند

یک تشبیه بهتر برای وزن شرکت‌های بزرگ فناوری، شرکت هند شرقی بریتانیا است که از سال ۱۶۰۰ تا اوایل دهه ۱۸۰۰، به عنوان یک نهاد مستقل با استقلال بسیار زیاد و با حمایت تاج و تخت بریتانیا عمل می‌کرد. شرکت هند شرقی بخش بزرگی از امپراتوری بریتانیا را ایجاد کرد، مستعمرات را با ارتش خود به دست آورد، اداره خود را مدیریت کرد و به تاج و تخت سهمی داد. حدود دو قرن طول کشید تا دولت-ملت دوباره خود را تثبیت کند و شرکت را منحل کند. به همین ترتیب، غول‌های فناوری امروزی قدرت عظیمی را جمع‌آوری کرده‌اند و از طریق کنترل زیرساخت‌ها و پلتفرم‌های دیجیتال که به بخش جدایی‌ناپذیر جامعه مدرن تبدیل شده‌اند، بر رفتار دولت تأثیر می‌گذارند. سوال اصلی این است: آیا دولت-ملت‌ها کنترل کافی را دوباره به دست خواهند آورد؟

احتمالاً، اما نه به این زودی‌ها. شرکت‌های بزرگ فناوری در حال شکل‌دهی به سیاست‌های دولت ایالات متحده هستند (که در طرح اقدام هوش مصنوعی که اخیراً توسط دولت منتشر شده، مشهود است) تا حدودی به این دلیل که کنگره اخیراً شروع به اقدام کرده است و تا حدی به دلیل سرعت بالای تغییرات. علاوه بر این، از زمانی که چین مدل هوش مصنوعی DeepSeek خود را عرضه کرد، ژئوپلیتیک همزیستی بین دولت و صنعت را تقویت کرده است. برخلاف چین، که شرکت‌های فناوری کاملاً تحت کنترل حزب-دولت هستند، در ایالات متحده، شرکت‌های بزرگ فناوری بر رویکرد دولت تأثیر می‌گذارند. ردپای سیاسی این صنعت در واشنگتن با افزایش نظارت عمومی گسترش یافته است: غول‌های فناوری ایالات متحده در سال ۲۰۲۴، ۶۱ میلیارد دلار برای لابی‌گری هزینه کردند. ماسک به تنهایی ۲۸۸ میلیون دلار در سال گذشته به کمپین‌های انتخاباتی کمک مالی کرد.

همزمان با اینکه تأثیر کامل اجتماعی، اقتصادی و ژئواستراتژیک هوش مصنوعی توسط انسان‌ها احساس می‌شود، سازندگان هوش مصنوعی به طور فزاینده‌ای ادعان می‌کنند که آنها به طور کامل نحوه عملکرد این سیستم‌ها را درک نمی‌کنند. بسیاری از سیستم‌ها در محیط‌های کم‌ریسک و مقاوم در برابر خطا عملکرد خوبی دارند. اما کنترل مهندسی در حوزه‌های پرتأثیر، مانند امور مالی، حقوق، مراقبت‌های بهداشتی و دفاع، که در آنها شکست‌ها عواقب جدی دارند، ضروری می‌شود. توهمات - پاسخ‌های نادرست یا ساختگی - در وظایف استدلال پیچیده رایج هستند و در برخی ارزیابی‌ها به ۸۰ درصد می‌رسند.

اگر مهندسانی که این سیستم‌ها را طراحی می‌کنند نتوانند به طور رضایت‌بخشی نحوه‌ی کار آنها را توضیح دهند، آیا واقعاً می‌توان هوش مصنوعی را کنترل کرد؟ مصطفی سلیمان، یکی از بنیانگذاران دیپ‌ماینند، مطمئن نیست. دیگران می‌گویند شناسایی حالت‌های خرابی دشوار است - و این باعث می‌شود که یک "کلید مرگ" (برای خاموش کردن هوش مصنوعی در صورت خرابی) غیرواقعی باشد. پوشه بنگیو، یکی از پیشگامان این حوزه، خواستار سرمایه‌گذاری در نظارت پیشرفته برای کنترل هوش مصنوعی عامل شده است، که می‌تواند به طور مستقل عمل کند، قبل از اینکه "چیزهایی بسازیم که بتوانند ما را نابود کنند".

همه موافق نیستند. یان لکان، دانشمند ارشد هوش مصنوعی متا، یکی دیگر از چهره‌های بنیادی، «تاامیدی هوش مصنوعی» را رد می‌کند. او معتقد است که مدل‌های فعلی از دستیابی به خودمختاری یا هوش عمومی فاصله زیادی دارند و هوش مصنوعی ایمن و قابل کنترل اساساً یک مشکل مهندسی است. با این حال، راه حل پیشنهادی او - نوآوری متن‌باز برای بهبود آزمایش و





اشکال‌زدایی - به چین دسترسی می‌دهد. در حال حاضر، مدل‌های اختصاصی هنوز بر اکوسیستم ایالات متحده تسلط دارند، اگرچه علاقه به مدل‌های متن‌باز رو به افزایش است.

بدون حفاظ، حکومتداری چندپاره

حتی با اینکه خطرات بیشتر قابل مشاهده و نگرانی‌ها گسترده‌تر شده‌اند، مقررات به سرعت پیشرفت نکرده‌اند. تلاش‌های بین‌المللی مانند اعلامیه بلجی که در سال ۲۰۲۳ به میزبانی بریتانیا برگزار شد، امیدی برای حاکمیت هماهنگ با قدرت‌های اصلی هوش مصنوعی ایجاد کرد. ایالات متحده، چین و اتحادیه اروپا به توسعه ایمن و انسان‌محور متعهد شدند. اما در اجلاس جهانی هوش مصنوعی پاریس ۲۰۲۵، هم ایالات متحده و هم بریتانیا از امضای اعلامیه بعدی در مورد اخلاق و ایمنی خودداری کردند. در واقع، جی. دی. ونس، معاون رئیس‌جمهور، علناً با مقررات «بیش از حد» ایمنی‌محور مخالفت کرد.

با این حال، واقعیت این است که علم برای ارزیابی عملکرد و ایمنی ناکافی است. این امر به ویژه نگران‌کننده است زیرا هوش مصنوعی از ابزاری برای تقویت قابلیت‌های انسانی - مانند هوش مصنوعی مولد (به ChatGPT فکر کنید) - به "عامل" تبدیل می‌شود و بدون نظارت انسان به صورت خودکار عمل می‌کند (به وسایل نقلیه بدون راننده فکر کنید). با وجود افزایش نمای ریسک مرتبط، تنها دو درصد از تحقیق و توسعه جهانی هوش مصنوعی به ایمنی و تنها پنج درصد به هماهنگی انسانی اختصاص داده شده است. بسیاری از محققان برجسته ایمنی با ناامیدی از شرکت‌های بزرگ فناوری خارج شده‌اند. در همین حال، استانداردهای جهانی الزام‌آور و مکانیسم‌های نظارتی قابل اجرا برای هوش مصنوعی عامل وجود ندارد. علاوه بر این، هیچ پروتکل فنی برای رسیدگی به عدم هماهنگی (زمانی که هوش مصنوعی اهداف مورد نظر طراحان خود را دنبال نمی‌کند) و همچنین مکانیسم‌های قابل اعتمادی برای مداخله ایمن در برابر خرابی (مانند یک کلید قطع) وجود ندارد.

اگرچه سرعت توسعه هوش مصنوعی، لزوم مقررات انعطاف‌پذیر را ایجاد می‌کند، اما به نظر می‌رسد ایالات متحده از این قاعده مستثنی است. چین و اتحادیه اروپا، با وجود داشتن سیستم‌های حاکمیتی بسیار متفاوت، بیشتر بر ایمنی تمرکز دارند. در این چشم‌انداز نظارتی پراکنده، هماهنگی جهانی هم دشوارتر و هم ضروری‌تر می‌شود.

هشدارهای نادیده گرفته شده

اکنون زمان آن رسیده است که یک قدم به عقب برداریم و دوباره به سوالات سطح اول بپردازیم: هوش مصنوعی بشریت را به کجا می‌برد؟ آیا جایگزین انسان‌ها خواهد شد؟ حتی آرمان‌شهرگرایان فناوری مانند سم آلتمن و ایلان ماسک نیز ابراز نگرانی کرده‌اند. ماسک تخمین می‌زند که ۲۰ درصد احتمال وجود دارد که هوش مصنوعی بتواند بشریت را نابود کند؛ آلتمن معتقد است که می‌تواند بر انسان‌ها به طور کامل غلبه کند. در مارس ۲۰۲۳، ماسک به همراه استیو ورنیاک، کارآفرین فناوری، و دها دانشمند هوش مصنوعی، نامه‌ای سرگشاده امضا کردند و خواستار توقف آموزش سیستم‌های پیشرفته‌تر هوش مصنوعی شدند و صریحاً پرسیدند: «آیا باید ذهن‌های غیرانسانی را توسعه دهیم که در نهایت ممکن است از نظر تعداد، هوش، منسوخ شدن و جایگزینی ما باشند؟» آلتمن این نامه را امضا نکرد، اما در آن به نقل از بیانیۀ OpenAI آمده است که پروژه‌های پیشرفته هوش مصنوعی باید سرعت افزایش منابع محاسباتی برای آموزش مدل را محدود کنند.

با این وجود، توسعه هوش مصنوعی با سرعت سرسام‌آوری در حال پیشرفت است. توسعه‌دهندگان درگیر رقابتی گيج‌کننده برای تصاحب آینده هستند که با اعتقاد به اجتناب‌ناپذیری هوش مصنوعی و ترس از دست دادن یک ثروت بادآورده مالی و ژئوپلیتیکی، به آن دامن زده می‌شود. شور و شوق غیرمنطقی بر نظم بازار غلبه کرده است. علیرغم بازده‌های به تعویق افتاده و خطرات حل نشده، سرمایه با شور و شوقی تقریباً مذهبی به توسعه هوش مصنوعی سرازیر می‌شود: بین سال‌های ۲۰۱۳ تا ۲۰۲۴، ۴۷۱ میلیارد دلار به سرمایه‌گذاری‌های هوش مصنوعی ایالات متحده سرازیر شده است و ۳۰۰ میلیارد دلار دیگر نیز تنها برای سال ۲۰۲۵ پیش‌بینی شده است.

با این حال، خطرات فرضی نیستند و هشدارها واضح هستند و توسط کسانی که خود سیستم‌ها را می‌سازند، صادر می‌شوند. پاسخگویی وجود ندارد، در قربانگاه شتاب قربانی می‌شود و فضای کمی برای نظارت یا اصلاح باقی می‌گذارد. بدون قوانین جهانی مشترک برای محافظت از نوآوری انسان‌محور، مدیریت هوش مصنوعی در معرض خطر شکل‌گیری صرفاً توسط قدرت و سود قرار

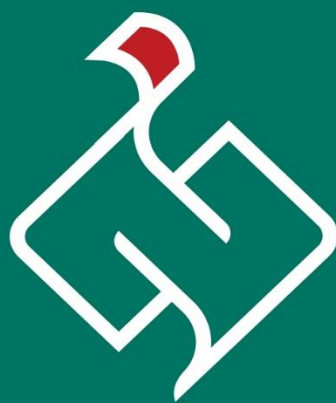


می‌گیرد. ادغام اخلاق و آینده‌نگری در توسعه فناوری، هر چه این پویایی بیشتر ادامه یابد، دشوارتر خواهد شد و احتمالاً عواقب عمیقی برای بشریت خواهد داشت.

آینده، بدون نظارت

به ندرت پیش آمده است که فناوری‌ای با چنین قدرت دگرگون‌کننده - و چنین خطرات و اختلالات قابل مشاهده و قابل درک - با چنین آمادگی کمی به کار گرفته شود. همانطور که داریو آمودی، مدیرعامل شرکت آنتروپیک، اخیراً هشدار داد: «شما نمی‌توانید جلوی قطار بایستید و آن را متوقف کنید.» اما می‌توانید - و باید - «آن را 10 درجه در جهتی متفاوت از جایی که می‌رفت هدایت کنید.» این چالش اصلی ایالات متحده در مسابقه هوش مصنوعی است.





انديشكده فرهنگ و توسعه

 WWW.CDSTT.IR

 [CDSTT.IR](https://www.instagram.com/cdstt.ir)

 [CDSTT](https://twitter.com/cdstt)