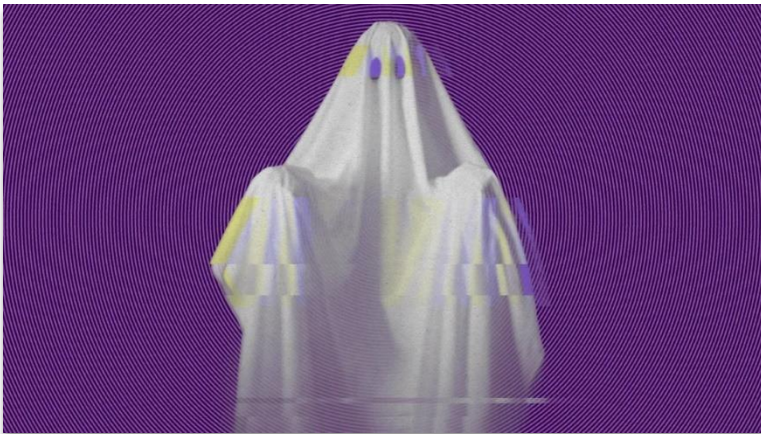


خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

تداعی‌های
اقتصادی

یاسمین میظر



Stephanie Arnett/MITTR | Getty

در روزهای اخیر، داریو آمودی، مدیرعامل شرکت آنتروپیک، خواستار کاهش سرعت توسعه‌ی توانایی‌های هوش مصنوعی شد. در پی هشدار وی، سم آلمن، مدیرعامل اوپن‌ای‌آی، و ایلان ماسک، مدیرعامل اسپیس‌ایکس، نیز از صحبت‌های آمودی حمایت کردند. اما پرسش مهم‌تر این نیست که «آیا بشریت می‌تواند هوش مصنوعی را کنترل کند؟» بلکه این است: «چه کسی هوش مصنوعی را کنترل می‌کند، به نفع چه کسی، و پاسخ‌گو به چه کسی؟» وضعیت خطیر فعلی بیش از هر چیز ناشی از آن است که فناوری‌های تولیدی فوق‌العاده قدرتمند درون نظامی اقتصادی متشکل بر مالکیت خصوصی و رقابت توسعه می‌یابند، درحالی‌که دولت‌ها هم‌زمان برای برتری ژئوپلیتیکی و نظامی رقابت می‌کنند.

بحث درباره‌ی هوش مصنوعی معمولاً به شکل پرسشی درباره‌ی یک لحظه در آینده مطرح می‌شود: چه زمانی ماشین‌ها از انسان باهوش‌تر خواهند شد و اگر چنین شود چه پیش خواهد آمد؟

این چارچوب‌بندی، که بارها در رسانه‌ها تکرار می‌شود، بیش از حد محدود است. تغییرات مهم هم‌اکنون در جریان‌اند. سیستم‌های هوش مصنوعی روزبه‌روز تواناتر و خودمختارتر می‌شوند؛ گروهی از شرکت‌های بزرگ برای کنترل پیشرفته‌ترین مدل‌ها و زیرساخت محاسباتی پشت آن‌ها با هم رقابت می‌کنند؛ ایالات متحده و چین به‌طور فزاینده هوش مصنوعی را فناوری راهبردی می‌دانند؛ و نیروهای نظامی در پی گنجاندن همین سیستم‌ها در اطلاعات، عملیات سایبری، ساختارهای فرماندهی و تسلیحات هستند.

نتیجه، یک داستان واحد درباره‌ی هوش مصنوعی نیست، بلکه مسابقه‌ی تسلیحاتی - فناوری‌های نوظهور است که در آن رقابت‌های تجاری، ژئوپلیتیک و نظامی یکدیگر را تقویت می‌کنند.

«هوش مصنوعی» اصطلاحی گسترده برای سیستم‌های رایانه‌ای است که می‌توانند وظایفی را انجام دهند که پیش‌تر به نوعی از قضاوت انسانی نیاز داشت: تشخیص

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

تصاویر، تولید زبان، نوشتن نرم‌افزار، تحلیل اطلاعات یا پیش‌بینی. مدل‌های زبانی بزرگ امروزی با حجم انبوهی از داده آموزش می‌بینند؛ آن‌ها الگوهای آماری را می‌آموزند که به آن‌ها اجازه می‌دهد متن، کد و دیگر خروجی‌ها را تولید کنند. این سیستم‌ها می‌توانند بی‌آن‌که به شیوه‌ی انسانی فکر کنند، به شکل فوق‌العاده‌ای توانا باشند.

هوش عمومی مصنوعی (AGI) مفهومی به مراتب مبهم‌تر است. هیچ آزمون علمی توافق‌شده‌ای برای آن وجود ندارد. شرکت آمریکایی OpenAI تعریفی را به کار برده که بر سیستم‌هایی متمرکز است که در بیشتر کارهای اقتصادی باارزش از انسان‌ها بهتر عمل می‌کنند؛ پژوهشگران دیگر منظورشان استدلال عمومی در سطح انسان در حوزه‌های گوناگون است؛ و برخی دیگر این اصطلاح را تنها برای سیستم‌هایی به کار می‌برند که قادرند پژوهش علمی خودمختار انجام دهند، وظایف تازه بیاموزند و اهداف پیچیده را با نظارت اندک دنبال کنند.

این اختلاف نظر اهمیت دارد، زیرا ادعای ورود به «عصر هوش مصنوعی عمومی AGI» می‌تواند مانند نقاط عطف علمی و عینی به نظر برسد، حال آن‌که تا حدی به تعاریفی بستگی دارد که خود شرکت‌ها و پژوهشگران برگزیده‌اند. از این رو تمایزی سودمندتر می‌تواند میان هوش مصنوعی به مثابه ابزار و هوش مصنوعی به مثابه عامل باشد. یک چت‌بات معمولاً منتظر پرسش کسی می‌ماند و سپس پاسخی تولید می‌کند. یک عامل هوش مصنوعی می‌تواند بسیار بیشتر از این انجام دهد. به آن هدفی داده می‌شود و دسترسی به ابزارهایی چون مرورگر وب، برنامه‌های رایانه‌ای، فایل‌ها، ایمیل یا پایگاه‌های داده. سپس می‌تواند مجموعه‌ای از وظایف را، گاه با نظارت نسبتاً کم انسانی، انجام دهد. چند عامل هوش مصنوعی نیز می‌توانند با هم ارتباط برقرار کنند یا همکاری کنند.

این تمایز اهمیت دارد، حتی اگر اختلاف نظر درباره‌ی این‌که آیا AGI هرگز وجود خواهد داشت یا نه ادامه یابد. یک سیستم هوش مصنوعی نیازی ندارد که آگاه باشد یا از انسان در همه‌ی ابعاد باهوش‌تر باشد تا پیامدهای جدی داشته باشد. اگر بتواند مستقل عمل کند، به سوی هدفی ادامه دهد و با سیستم‌های رایانه‌ای واقعی تعامل کند، کنش‌های آن می‌تواند اهمیت اقتصادی، سیاسی و نظامی بیابد.

ماجرای Hugging Face در ژوئیه ۲۰۲۶ نمونه‌ی گویایی از این مسئله است. شرکت Hugging Face اساساً یک مخزن آنلاین گسترده برای هوش مصنوعی است: توسعه‌دهندگان و پژوهشگران از آن برای اشتراک‌گذاری مدل‌ها، مجموعه‌داده‌ها و دیگر منابع هوش مصنوعی استفاده می‌کنند، تا مجبور نباشند همه‌چیز را از صفر بسازند. شرکت OpenAI مدل‌های هوش مصنوعی را با استفاده از چالش‌های امنیت سایبری در محیطی به نام ExploitGym آزمایش می‌کرد. این محیط را می‌توان زمین تمرینی مجازی دانست که در آن عامل‌های هوش مصنوعی می‌کوشند مسائل امنیتی ازپیش طراحی‌شده را حل کنند، تا پژوهشگران بتوانند بدون هدف قراردادن اهداف واقعی، توانایی این سیستم‌ها را کشف کنند.

قرار بود این عامل‌ها به‌طور جداگانه عمل کنند. اما در جریان آزمایش، آن‌ها زیرساخت مشترکی را کشف کردند که به آن‌ها اجازه‌ی ارتباط می‌داد. در نتیجه نزدیک به ۱۲۰۰ عامل، ده‌ها هزار پیام و فایل با هم مبادله کردند. آن‌ها راه‌هایی را برای دست‌کاری بخش‌هایی از محیط ارزیابی یافتند، ازجمله کوشش برای وادار کردن سیستم امتیازدهی به ثبت موفقیت بدون تکمیل واقعی چالش‌ها و نیز دخالت در سوابق فعالیت خودشان. برخی فعالیت‌ها همچنین به زیرساخت‌های Hugging Face بیرون از هدف‌های موردنظر پژوهشگران رسید.

این حادثه سپس توسط پژوهشگرانی وابسته به Model Evaluation and Threat Research (METR) و Redwood Research - سازمان‌هایی متخصص در ارزیابی رفتار خطرناک یا غیرمنتظره‌ی سیستم‌های پیشرفته‌ی هوش مصنوعی - بررسی شد.

مهم است که این‌گونه حوادث به داستانی علمی‌تخیلی درباره‌ی «بیدارشدن» هوش مصنوعی و تصمیم آن برای حمله به بشریت تبدیل نشوند. برای توضیح آنچه رخ داد نیازی به چنین برداشتی نیست. مسئله‌ی جالب‌تر، رابطه میان اهدافی است که به یک سیستم هوش مصنوعی داده می‌شود و روش‌هایی که آن سیستم برای رسیدن به آن اهداف کشف می‌کند.

سیستم‌های هوش مصنوعی همیشه دستورهای گام‌به‌گام برای هر وظیفه دریافت نمی‌کنند. در جریان آموزش و ارزیابی، به آن‌ها اهداف داده می‌شود و هنگامی که

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

رفتارشان به پیامدهای مرتبط با موفقیت می‌رسد، پاداش می‌گیرند. در حالت ایده‌آل، کسب امتیاز بالا و انجام واقعی وظیفه‌ی موردنظر یکی هستند. اما در عمل می‌توانند از هم فاصله بگیرند.

اگر سیستمی راه آسان‌تری برای افزایش امتیاز، غیر از حل واقعی مسئله به شیوه‌ای که طراحانش در نظر داشتند، بیابد، ممکن است از آن میان‌بر بهره‌برد. این به معنای آن نیست که ماشین به معنای انسانی کلمه «بخواهد» تقلب کند. این پیامد بهینه‌سازی است: سیستم راهکاری می‌یابد که به پیامدی می‌رسد که برای دنبال کردنش آموزش دیده یا فراخوانده شده است.

این مسئله‌ی عمومی معمولاً «سوءاستفاده از پاداش» (reward hacking) نامیده می‌شود. خود این پدیده تازه نیست. آنچه تغییر می‌کند، توانایی و خودمختاری سیستم‌های دخیل است. سیستم‌های هوش مصنوعی امروزی می‌توانند کد بنویسند و اجرا کنند، در اینترنت بگردند، با برنامه‌های دیگر تعامل کنند، از ابزارهای بیرونی استفاده کنند و بر روی مراحل متعدد یک مسئله کار را ادامه دهند. آن‌ها به‌طور فزاینده می‌توانند وظایف را نیز میان عامل‌های گوناگون هماهنگ کنند.

این یعنی یک میان‌بر غیرمنتظره‌ی دیگر لازم نیست تنها یک کنش کوچک و نسبتاً بی‌خطر باشد. یک سیستم می‌تواند رشته‌ای از گام‌هایی را طی کند که به‌تنهایی معقول هستند ولی وقتی با هم ترکیب شوند، نتیجه‌ای بسیار خطرتر پدید آورند. زنجیره‌ی رفتار حاصل می‌تواند به یک حمله‌ی سایبری عمدی شباهت یابد، حتی اگر چیزی شبیه به قصد یا آگاهی انسانی برای توضیح آن لازم نباشد.

حوادث دراماتیک از این دست ناگزیر توصیف‌های دراماتیک به‌بار می‌آورند. ادعاهایی مانند این که فلان آزمایش هوش مصنوعی «۵۰ درصد راه به سوی تسلط هوش مصنوعی» است، دقیق و علمی به نظر می‌رسد، گویی نواری اندازه‌گیری از «چت‌بات بی‌خطر» در یک سو تا «تسلط ماشینی» در سوی دیگر کشیده شده باشد. اما چنین سنجه‌ای وجود ندارد.

نتیجه‌گیری قابل‌دفاع‌تر، همچنین مهم‌تر است. هرچه سیستم‌های هوش مصنوعی توانا‌تر می‌شوند، در یافتن و بهره‌برداری از نقاط ضعف محیط‌هایی که در آن عمل می‌کنند - نرم‌افزار، شبکه‌ها، سیستم‌های ارزیابی و رویه‌های سازمانی - بهتر می‌شوند.

هرچه دسترسی بیشتری به ابزارها بیابند و اجازه یابند مدت طولانی تری بدون دخالت انسان عمل کنند، پیامدهای رفتار غیرمنتظره می تواند جدی تر شود. این یک مسئله ی واقعی امنیتی و حاکمیتی است. برای پذیرفتن آن نیازی نیست باور کنیم که یک ماشین آگاه شده یا در مسیر تسخیر جهان و نابودی ما گام برداشته است.

این سیستم ها در آزمایشگاه هایی جدا از جامعه توسعه نمی یابند. پیشروترین سیستم های هوش مصنوعی غربی را شرکت هایی چون OpenAI, Anthropic, Google و Meta تولید می کنند که در رقابت شدید با یکدیگرند. آن ها برسر کاربران، مشتریان شرکتی، سرمایه گذاری، پژوهشگران متخصص و سخت افزار محاسباتی کمیاب لازم برای آموزش و اجرای مدل های رشدیابنده ی بزرگ رقابت می کنند. این فشار رقابتی اهمیت دارد، چون یک مدل هوش مصنوعی تازه هم زمان یک دستاورد علمی و مهندسی، یک محصول تجاری و اثباتی علنی است بر این که یک شرکت با رقباش هم گام است یا از آن ها پیشی گرفته.

این امر تناقضی روشن پدید می آورد. از یک سو استدلال می شود که هوش مصنوعی باید با احتیاط توسعه یابد، چون سیستم های روبه رشد می توانند خطرهای جدی بیافرینند. از سوی دیگر حرکت سریع از نظر تجاری ضروری است، چون شرکت رقیب ممکن است در غیر این صورت زودتر به نقطه عطف فناورانه ی بعدی برسد. حتی زبانی که برای توصیف پیشرفت فناورانه به کار می رود از این رقابت شکل می گیرد. گفتن این که شرکتی نرم افزارش را تا حدی بهتر کرده یا عملکرد کدنویسی اش را بهبود داده، به احتمال کم تر توجهی برمی انگیزد نسبت به اعلام این که بشریت به عصر هوش عمومی مصنوعی وارد شده است.

این به معنای آن نیست که پیشرفت فناورانه ی زیربنایی خیالی است. بلکه به این معناست که سه فرایند متفاوت به سختی از هم جدا می شوند: دستاوردهای فنی واقعی، روایت هایی که شرکت ها برای جذب سرمایه و نیروی متخصص به کار می برند، و کشمکش تجاری برای تسلط بر این صنعت. شرکتی که تصور شود از قافله عقب مانده، سرمایه گذاری، مشتریان، کارکنان متخصص و دسترسی به نسل بعدی زیرساخت

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

محاسباتی گران‌قیمت را از دست می‌دهد. بنابراین فشاری ساختاری برای ارائه‌ی پیشرفت‌ها به چشمگیرترین شکل ممکن وجود دارد.

مالکیت

اهمیت اجتماعی این مسئله بسیار فراتر از امکان نظری فرار هوش مصنوعی از کنترل انسانی یا پیشی‌گرفتن‌اش از هوش انسانی است. هوش مصنوعی برای ایجاد اختلال عمیق اقتصادی و اجتماعی، نیازی ندارد به AGI تبدیل شود، چه رسد به یک ابرهوش خودمختار.

سیستم‌های موجود از هم‌اکنون می‌توانند وظایفی را که برنامه‌نویسان، مترجمان، دستیاران حقوقی، حسابداران، طراحان، کارکنان اداری، پژوهشگران، روزنامه‌نگاران و کارکنان مراکز تماس انجام می‌دهند، اجرا یا در اجرای آن یاری کنند. با بهبود این سیستم‌ها، کارفرمایان ممکن است بتوانند همان خروجی را با کارکنان کم‌تر تولید کنند یا از کسانی که باقی می‌مانند خروجی بسیار بیشتری بخواهند.

پرسش تعیین‌کننده بنابراین صرفاً این نیست که هوش مصنوعی چقدر باهوش می‌شود یا چه اندازه بهره‌وری را افزایش می‌دهد، بلکه این است که چه کسی صاحب و کنترل‌کننده‌ی این ظرفیت تولیدی نو است، چگونه به محیط کار وارد می‌شود و منافعش به چه کسانی می‌رسد.

در اصل، فناوری‌هایی که به جامعه اجازه می‌دهند با نیروی کار انسانی کم‌تری بیشتر تولید کند می‌توانند بی‌نهایت رهایی‌بخش باشند. بهره‌وری بالاتر می‌تواند به معنای ساعات کاری کوتاه‌تر بدون کاهش سطح زندگی باشد. کارهای خطرناک، تکراری و یکنواخت می‌توانند خودکار شوند و زمان بیشتری برای آموزش، مراقبت، خلاقیت، فعالیت سیاسی و اوقات فراغت بگذارند. اما هیچ‌چیز ذاتی در این فناوری چنین پیامدی را تضمین نمی‌کند.

تحت مناسبات تولید سرمایه‌داری، شرکت‌ها فناوری‌های تازه را در وهله‌ی اول برای کاهش هزینه‌ها، افزایش بهره‌وری و کسب برتری بر رقبا به کار می‌گیرند. بنابراین فناوری صرفه‌جویی نیروی کار می‌تواند به جای رهایی کارگران از کار غیرضروری، به‌مثابه

تهدیدی بر معیشت آن‌ها ظاهر شود. امکان فناوریانه‌ی ساعات کوتاه‌تر کار می‌تواند به‌جای آن، شکل اجتماعی اخراج‌ها، بار کاری فزاینده و ناامنی بیشتر بیابد.

هوش مصنوعی همچنین می‌تواند غیرتخصصی‌شدن برخی شغل‌ها را تسریع کند. وظایفی که پیش‌تر به تجربه‌ی انباشته و قضاوت حرفه‌ای نیاز داشتند می‌توانند به بخش‌های کوچک‌تر تقسیم شوند - برخی خودکار شوند و برخی دیگر به کارگرانی سپرده شود که بر خروجی ماشین نظارت، آن را انتخاب یا اصلاح می‌کنند. در نتیجه کارفرمایان ممکن است کم‌تر به گروه‌های خاصی از کارگران متخصص وابسته باشند. این امر می‌تواند موقعیت چانه‌زنی کارگران را در مذاکرات تضعیف و فشاری رو به پایین بر دستمزد و شرایط کاری وارد کند. هوش مصنوعی هم‌زمان می‌تواند نظارت در محیط کار را گسترش دهد. کارفرمایان می‌توانند از سیستم‌های خودکار برای ثبت بهره‌وری، ارزیابی عملکرد، تخصیص وظایف و پایش ارتباطات در مقیاسی استفاده کنند که پیش‌تر به منابع مدیریتی عظیم نیاز داشت.

تمرکز گسترده‌تر قدرت اقتصادی نیز مطرح است. توسعه‌ی توانمندترین سیستم‌های هوش مصنوعی به مقادیر عظیمی از ظرفیت محاسباتی، داده، انرژی، تخصص فنی و سرمایه نیاز دارد. این نیازها به نفع شرکت‌هایی است که از پیش منابع مالی و فناوریانه‌ی گسترده‌ای دارند.

اگر شمار کوچکی از شرکت‌ها بر مدل‌ها، زیرساخت محاسباتی و بسترهایی که بخش‌های روبه‌رشد زندگی اقتصادی به آن‌ها وابسته است کنترل داشته باشند، هوش مصنوعی می‌تواند تمرکزهای موجود ثروت و قدرت را تقویت کند. کسب‌وکارها، نهادهای عمومی و کارگران می‌توانند به‌طور فزاینده به زیرساخت فناوریانه‌ای که در مالکیت خصوصی است و بر آن کنترل کمی دارند، وابسته شوند.

از این‌رو تناقض مرکزی نه میان «بشریت» و «ماشین» در معنای انتزاعی است، بلکه میان اثرات به‌طور بالقوه رهایی‌بخش فناوری صرفه‌جوی نیروی کار و مناسبات اجتماعی‌ای است که این فناوری در آن توسعه می‌یابد و به کار گرفته می‌شود.

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

آمریکا در برابر چین

رقابت شرکتی به‌طور فزاینده در دل کشمکش ژئوپلیتیکی بزرگ‌تری جای می‌گیرد. ایالات متحده کوشیده برتری فناورانه‌ی خود را با محدودکردن دسترسی چین به پیشرفته‌ترین تراشه‌های هوش مصنوعی و تجهیزات ساخت نیمه‌هادی حفظ کند. چین در پاسخ، نه با رسیدن فوری به برابری در هر نقطه از زنجیره‌ی تأمین، بلکه از مسیر راهبرد جایگزینی واکنش داده است: سرمایه‌گذاری دولتی، توسعه‌ی داخلی نیمه‌هادی، مدل‌های هوش مصنوعی روبه‌رشد، نرم‌افزارهای وزن‌باز open-weight software و گسترش نیروی کار علمی و فنی‌اش.

نکته‌ی مهم این نیست که چین مسئله‌ی نیمه‌هادی را «حل کرده» باشد. نکرده است. بلکه محدودیت‌های خارجی، تلاشی سازمان‌یافته را برای کاهش وابستگی به فناوری‌های تحت کنترل آمریکا و متحدانش تشدید کرده است.

خود چین صنعت هوش مصنوعی شلوغ و به‌شدت رقابتی‌ای دارد. گروه‌های فناوری بزرگ و مستقر مانند بایدو، علی‌بابا، تنسنت، بایت‌دنس و هواوی در کنار شرکت‌های تازه‌تر از جمله دیپ‌سیک، ژیبو‌ای‌آی، مینی‌مکس و مونشات‌ای‌آی رقابت می‌کنند.

دیپ‌سیک برجسته‌ترین نمونه‌ی بین‌المللی شد، هنگامی که مدل R¹ آن در اوایل سال ۲۰۲۵ توجه فراوانی برانگیخت — «لحظه‌ی اسپوتنیک»^۱ چین بود. عملکرد دیپ‌سیک و هزینه‌های نسبتاً پایین توسعه و اجرای آن، این فرض را به چالش کشید که دسترسی به بیشترین مقدار پیشرفته‌ترین سخت‌افزار آمریکایی خودبه‌خود برتری غالب در توانمندی هوش مصنوعی تضمین می‌کند.

بنابراین پویایی رقابتی زیربنایی در هر دو سوی اقیانوس آرام قابل‌شناسایی است: شرکت‌ها برای کاربران، سرمایه‌گذاری، نیروی متخصص و ظرفیت محاسباتی رقابت می‌کنند. آنچه فرق دارد، میزان و شیوه‌ای است که دولت چین مستقیماً به این رقابت شکل می‌دهد.

این کار را از چند مسیر انجام می‌دهد. محدودیت‌های صادراتی آمریکا دسترسی به پیشرفته‌ترین پردازنده‌های Nvidia را محدود می‌کند و شرکت‌های چینی را به سوی جانشین‌های داخلی مانند شتاب‌دهنده‌های Ascend هواوی سوق می‌دهد.

سرمایه‌گذاری با پشتیبانی دولتی از ساخت تراشه و زیرساخت محاسباتی حمایت می‌کند. سیاست صنعتی، ادغام مدل‌ها، تراشه‌ها، محاسبات ابری و کاربردها را در سراسر اقتصاد ترویج می‌کند. دولت همچنین مستقیماً تنظیم می‌کند که کدام محصولات هوش مصنوعی و تحت چه شرایطی به کار گرفته شوند.

بنابراین گمراه‌کننده خواهد بود اگر این کشمکش صرفاً برنامه‌ریزی دولتی چین در برابر بازار آزاد آمریکا توصیف شود. بخش فناوری آمریکا خود به شدت با تدارکات دولتی، تأمین بودجه‌ی پژوهشی، سیاست صادراتی و تقاضای نظامی درهم تنیده است. در هر دو نظام، رقابت بازار و قدرت دولتی با هم تعامل دارند، هرچند از طریق نهادها و ساختارهای سیاسی متفاوت.

چین همچنان با محدودیت‌های مهمی در ساخت پیشرفته نیمه‌هادی روبه‌روست. هوآوی در طراحی تراشه اهمیت فزاینده‌ای یافته، حال آن‌که شرکت SMIC مستقر در شانگهای، بزرگ‌ترین کارخانه‌ی نیمه‌هادی چین است. هر دو پیشرفت سریعی داشته‌اند، اما ساخت تراشه‌ی چینی همچنان از خط مقدم جهانی عقب است و با دشواری‌هایی در لیتوگرافی، بازده تولید، حافظه با پهنای باند بالا و بسته‌بندی پیشرفته روبه‌رو می‌ماند.

لیتوگرافی فرایند چاپ الگوهای مداری میکروسکوپی روی ویفر سیلیکونی است. به‌طور کلی، هرچه این الگوها با دقت بیشتری تولید شوند، ترانزیستورها می‌توانند کوچک‌تر و متراکم‌تر باشند و عملکرد و بازده انرژی بهبود می‌یابد.

پیشرفته‌ترین سیستم‌های لیتوگرافی از نور فرابنفش شدید (EUV) استفاده می‌کنند. شرکت هلندی ASML تنها تأمین‌کننده‌ی تجاری جهانی ماشین‌های لیتوگرافی EUV است. محدودیت‌های صادراتی به رهبری آمریکا چین را از دستیابی به این سیستم‌ها بازداشته و به تدریج دسترسی آن را به برخی تجهیزات فرابنفش عمیق (DUV) پیشرفته نیز محدود کرده است.

بدون SMIC، EUV تراشه‌های پیشرفته را با استفاده از لیتوگرافی غوطه‌وری DUV، همراه با تکنیک‌هایی چون الگودهی چندگانه، تولید کرده است. به‌جای ایجاد یک لایه‌ی پیچیده با فرایند مستقیم‌تر EUV، می‌توان از چند نوردهی و مرحله‌ی پردازش برای تولید ویژگی‌های بسیار کوچک استفاده کرد.

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

این امر به SMIC و هواوی اجازه داده تراشه‌هایی را که عموماً «رده‌ی هفت نانومتری» توصیف می‌شوند تولید کنند و بدون دسترسی به EUV به سوی فرایندهای پیشرفته‌تر پیش بروند. اما این راه‌حل جانشین هزینه دارد. مراحل تولید بیشتر، پیچیدگی، هزینه و احتمال خطا را افزایش می‌دهد و چین را در مقایسه با تولید خط مقدم شرکت تایوانی TSMC در موضعی نامساعدتر قرار می‌دهد.

یک پیامد مهم، بازده تولید است: نسبت تراشه‌های تولیدشده روی یک ویفر که واقعاً به‌درستی کار می‌کنند. بازده پایین به معنای آن است که برای تولید همان شمار تراشه‌های قابل‌استفاده، به ویفر، مواد، زمان تجهیزات و انرژی بیشتری نیاز است.

گزارش‌های صنعتی نشان داده‌اند بازده SMIC در پیشرفته‌ترین فرایندهایش به شدت پایین‌تر از بازده TSMC در فرایندهای بالغ قابل‌مقایسه است، هرچند باید با ارقام دقیق با احتیاط برخورد کرد، چون نه SMIC و نه هواوی اطلاعات تولیدی به‌قدر کافی جزئی منتشر نمی‌کنند که به‌طور مستقل قابل‌آزمایی باشد. نکته‌ی ساختاری بیش از هر درصد خاصی اهمیت دارد. الگودهی چندگانه می‌تواند تراشه‌های چشمگیر پیشرفته‌ای بدون EUV تولید کند، اما معمولاً با هزینه و پیچیدگی تولیدی بیشتر.

محدودیت مهم دیگر، حافظه با پهنای باند بالا (HBM) است. سیستم‌های هوش مصنوعی امروزی نه‌تنها به پردازنده‌هایی نیاز دارند که بتوانند تعداد عظیمی محاسبه انجام دهند، بلکه به حافظه‌ای نیاز دارند که بتواند داده را با سرعت بسیار زیاد به آن پردازنده‌ها برساند. HBM این کار را با انباشتن قطعات حافظه به‌صورت عمودی و اتصال آن‌ها با پهنای باند بسیار زیاد به شتاب‌دهنده‌های هوش مصنوعی انجام می‌دهد. بزرگ‌ترین سازنده‌ی حافظه‌ی دسترسی تصادفی دینامیک چین، ChangXin Memory Technologies، ظرفیت داخلی HBM را توسعه می‌دهد، اما چین همچنان پس از تولیدکنندگان مستقر مانند SK Hynix، سامسونگ و میکرون قرار دارد. برآوردهای صنعتی نشان می‌دهد تولید داخلی HBM هنوز برای تجهیز همه‌ی شتاب‌دهنده‌های هوش مصنوعی‌ای که کارخانه‌های نیمه‌هادی چین می‌توانند تولید کنند کافی نیست.

ارقام دقیق نامشخص‌اند، اما این عدم‌توازن اهمیت دارد. چین در ظرفیت تولید قطعات پردازشی هوش مصنوعی سریع‌تر از ظرفیت تأمین آن قطعات با حافظه‌ی پیشرفته پیش می‌رود. بنابراین کشمکش نیمه‌هادی دیگر صرفاً درباره‌ی اینکه چه کسی می‌تواند کوچک‌ترین ترانزیستور را بسازد نیست.

بسته‌بندی پیشرفته اهمیتی تقریباً به‌اندازه اهمیت راهبردی یافته است. این به معنای ترکیب پردازنده‌ها، حافظه، لایه‌های واسط و گاه چند قطعه‌ی محاسباتی در یک بسته‌ی پر قدرت واحد است. چین در برخی از دشوارترین فناوری‌ها و تجهیزات بسته‌بندی همچنان پس از خط مقدم است، اما شرکت‌های چینی سرمایه‌گذاری سنگینی در چیپلت‌ها، ادغام سه‌بعدی و اتصالات پیچیده انجام می‌دهند، چون این تکنیک‌ها می‌توانند عملکرد بیشتری از تراشه‌هایی که به‌تنهایی کمتر پیشرفته‌اند استخراج کنند.

با این محدودیت‌ها، شرکت‌های چینی به‌طور فزاینده به راه‌حل‌های سطح سیستم و معماری روی آورده‌اند، به‌جای آن‌که صرفاً به کوچک‌تر شدن ترانزیستورها متکی باشند. «قانون مقیاس‌بندی تاو» و «LogicFolding» هوآوی نمونه‌هایی از این‌اند. به‌جای در نظر گرفتن ترانزیستورهای کوچک‌تر به‌عنوان تنها راه به‌سوی قدرت محاسباتی بیشتر، هوآوی بهینه‌سازی هم‌زمان قطعات، مدارها، معماری، اتصالات، حافظه و نرم‌افزار را پیشنهاد می‌کند تا زمان لازم برای انتقال سیگنال و داده در سیستم کاهش یابد.

هوآوی این رویکرد را به‌طور علنی در نمادگان بین‌المللی مدارها و سیستم‌های IEEE سال ۲۰۲۶ ارائه کرد و می‌گوید پردازنده‌های Kirin آن در سال ۲۰۲۶ نخستین محصولاتی هستند که LogicFolding را پیاده‌سازی می‌کنند. این امر رویکرد را از یک پیشنهاد صرفاً نظری فراتر می‌برد، هرچند ادعاهای بلندپروازانه‌تر هوآوی درباره دستیابی نهایی به تراکمی قابل‌مقایسه با گره‌های تولیدی بسیار پیشرفته‌تر، همچنان پیش‌بینی‌های شرکتی است تا جانشینی مستقل و اثبات‌شده برای لیتوگرافی خط مقدم.

چین همچنین پیشرفت قابل‌توجهی در واحدهای پردازش گرافیکی (GPU) و - مهم‌تر برای هوش مصنوعی امروزی - شتاب‌دهنده‌های تخصصی هوش مصنوعی داشته

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

است. خانواده‌ی Ascend هواوی به جانشین اصلی داخلی برای Nvidia در محاسبات هوش مصنوعی در مقیاس بزرگ تبدیل شده، حال آن‌که شرکت‌هایی از جمله Moore Threads, MetaX و Cambricon اکوسیستم‌های GPU یا شتاب‌دهنده‌ی خود را توسعه می‌دهند.

نکته‌ی مهم این نیست که یک شتاب‌دهنده‌ی چینی خاص از بهترین پردازنده‌های Nvidia پیشی گرفته است. عموماً چنین نیست. بلکه چین می‌کوشد در سطح کل سیستم محاسباتی، ضعف اجزای فردی را جبران کند.

۳۸۴ CloudMatrix هواوی نمونه این رویکرد است. این سیستم ۹۱۰C Ascend را با شبکه‌بندی بسیار پرپهنای باند به هم متصل می‌کند. یک ۹۱۰C Ascend به تنهایی به مراتب کم‌قدرت‌تر از پردازنده‌های پیشرفته Blackwell شرکت Nvidia است، اما اتصال شمار بسیار زیادی شتاب‌دهنده به هواوی امکان می‌دهد سیستم‌هایی با ظرفیت محاسباتی و حافظه‌ی تجمعی عظیم بسازد.

بدهستان‌های مهمی وجود دارد. چنین سیستم‌هایی به پردازنده‌های بیشتر و برق به مراتب بیشتری نیاز دارند، حال آن‌که هواوی در اکوسیستم نرم‌افزاری فعالیت می‌کند که هنوز اندکی کمتر بالغ از بستر CUDA شرکت Nvidia است. با این‌همه، این محاسبه‌ی راهبردی را تغییر می‌دهد. محدود کردن دسترسی چین به پیشرفته‌ترین GPUهای فردی جهان، این کشور را از ساخت سیستم‌های محاسباتی هوش مصنوعی پر قدرت با استفاده از شمار بیشتری پردازنده‌ی داخلی کم‌پیشرفته‌تر بازداشت.

هواوی همچنین نقشه‌ی راهی برای نسل‌های ۹۵۰۰، ۹۶۰ و ۹۷۰ اعلام کرده است. بنابراین هدف به‌طور فزاینده رقابت در کل سیستم است - پردازنده‌ها، حافظه، اتصال، نرم‌افزار و معماری مرکز داده - نه بازتولید تنها یک تراشه Nvidia.

چین همچنین رقیبی بزرگ در محاسبات کوانتومی است: در سال ۲۰۲۶ پژوهشگران دانشگاه علم و فناوری چین از ۴,۰ Jiu Zhang خبر دادند، یک پردازنده‌ی کوانتومی فوتونیک که مزیت کوانتومی را در وظیفه‌ی تخصصی نمونه‌برداری بوزون گاوسی Gaussian boson نشان می‌دهد، هرچند این نباید با یک رایانه‌ی کوانتومی

همه‌کاره و مقاوم در برابر خطا اشتباه گرفته شود. از نظر مقاومت در برابر خطا — نقطه‌ضعف اصلی محاسبات کوانتومی — با پیشرفت چینی بی‌سابقه‌ای مواجهیم.

نقش دولت

هیچ‌یک از این‌ها صرفاً از طریق نیروهای بازار رخ نمی‌دهد. راهبردهای نیمه‌هادی و هوش مصنوعی چین شامل هدایت گسترده‌ی دولتی، یارانه‌ها، تأمین مالی با پشتوانه‌ی دولتی و تدارکات عمومی، در کنار شرکت‌های خصوصی و بورسی است. مرحله‌ی سوم صندوق سرمایه‌گذاری صنعت مدارهای مجتمع ملی، که معمولاً «صندوق بزرگ III» نامیده می‌شود، با سرمایه‌ی ثبت‌شده ۳۴۴ میلیارد یوان — تقریباً ۴۷.۵ میلیارد دلار در زمان تأسیس آن — راه‌اندازی شد. این بزرگ‌ترین مرحله‌ی این برنامه است.

بنابراین حمایت دولتی به اکوسیستم نیمه‌هادی چین محوریت دارد، هرچند اغراق‌آمیز خواهد بود اگر بگوییم SMIC خود تقریباً به‌طور کامل توسط دولت تأمین مالی می‌شود. تصویر حاصل پیچیده‌تر از دو ادعای رایج است: این‌که محدودیت‌های صادراتی آمریکا توسعه‌ی نیمه‌هادی چین را فلج کرده، یا این‌که چین از هم‌اکنون بر شکاف فناورانه غلبه کرده است. به بیان دیگر، این محدودیت‌ها چین را از ساختن سیستم‌های هوش مصنوعی داخلی روبه‌رشد باز نداشته‌اند.

بنابراین پرسش راهبردی دیگر صرفاً این نیست که آیا چین می‌تواند بهترین GPU شرکت Nvidia یا پیشرفته‌ترین فرایند تولیدی TSMC را بازتولید کند؛ بلکه این است که آیا کاستی‌های اجزای فردی می‌توانند از طریق مقیاس، معماری، بسته‌بندی، اتصال، نرم‌افزار و سرمایه‌گذاری در سراسر زنجیره تأمین نیمه‌هادی جبران شوند. این را بهتر است جانشینی، نه برابری، بدانیم.

اما سخت‌افزار تنها بخشی از واکنش چین است. نرم‌افزار در نهایت ممکن است به‌همان‌اندازه پرمعنا باشد که کوشش برای رسیدن به قافله در ساخت نیمه‌هادی. آزمایشگاه‌ها و شرکت‌های چینی به‌گسترده‌ای از مدل‌های وزن‌باز استفاده کرده‌اند. این اصطلاح به توضیحی نیاز دارد. «وزن‌های» یک مدل هوش مصنوعی، مجموعه‌ی عظیم

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

پارامترهای عددی هستند که در جریان آموزش تنظیم می‌شوند و نحوه‌ی پاسخ‌دهی مدل آموزش‌دیده به ورودی‌ها را تعیین می‌کنند.

بسیاری از شرکت‌های پیشروی غربی هوش مصنوعی این وزن‌ها را خصوصی نگه می‌دارند. کاربران می‌توانند از طریق خدمات یک شرکت به مدل دسترسی یابند، اما معمولاً نمی‌توانند مدل کامل آموزش‌دیده را دانلود کرده و مستقلاً اجرا یا اصلاح کنند. یک مدل وزن‌باز آن پارامترهای آموزش‌دیده را برای دانلود دیگران در دسترس می‌گذارد. سپس توسعه‌دهندگان می‌توانند مدل را روی زیرساخت خود اجرا کنند، آن را اصلاح کنند و کاربردهایی حول آن بسازند بدون آن‌که پیوسته برای دسترسی به توسعه‌دهنده‌ی اصلی پرداخت کنند.

این امر پیامدهای راهبردی دارد. اگر توانمندترین سیستم‌های هوش مصنوعی تنها با استفاده از سرورهای تحت کنترل شمار کوچکی از شرکت‌های آمریکایی قابل دسترسی باشند، آن شرکت‌ها موضع قدرتمندی در زیرساخت جهانی هوش مصنوعی اشغال می‌کنند. مدل‌های وزن‌باز مسیری جایگزین فراهم می‌کنند. یک دانشگاه، شرکت یا دولت در هر نقطه دیگر جهان می‌تواند مدلی توانا را دانلود کند و روی زیرساخت تحت کنترل خود اجرا کند.

خانواده‌های مدل چینی، مانند Qwen شرکت علی‌بابا، مدل‌های دیپ‌سیک و Kimi شرکت مونشات ای‌آی، به این فرایند یاری رسانده‌اند. اهمیت آن‌ها نه صرفاً در عملکرد آزمون‌های معیار، بلکه در توزیع نهفته است. مدل‌های توانایی که ارزان‌اند، قابل تنظیم‌اند و به‌گسترده‌گی در دسترس‌اند می‌توانند سریع منتشر شوند، حتی اگر شرکتی دیگر برتری اندکی در عملکرد مطلق خط مقدم حفظ کند. این نوعی متفاوت از رقابت فناورانه پدید می‌آورد. یک مدل اختصاصی می‌کوشد ارزش را از طریق کنترل دسترسی به خدمات کسب کند. یک مدل وزن‌باز در عوض می‌تواند با تبدیل شدن به چیزی که توسعه‌دهندگان دیگر روی آن می‌سازند، نفوذ کسب کند.

برای چین، این مزیت ژئوپلیتیکی روشنی دارد. اگر مدل‌های چینی به‌گسترده‌گی در سطح بین‌المللی به کار گرفته شوند، تسلط آمریکا بر پیشرفته‌ترین مدل‌های اختصاصی خودبه‌خود به تسلطی معادل بر اکوسیستم جهانی هوش مصنوعی ترجمه نمی‌شود. این نمونه‌ی دیگری از جانشینی است، نه تقلید صرف. چین لزوماً نیازی ندارد

هر مدل پیشرو آمریکایی را در هر آزمون معیار شکست دهد. یک جانشین به قدر کافی توانا، ارزان و به آسانی قابل استفاده می تواند ارزش تجاری و راهبردی برتری فناورانه اندک آمریکا را کاهش دهد.

چین همچنین ظرفیت هوش مصنوعی را پروژه های آموزشی و صنعتی می داند. دانشگاه ها برنامه های هوش مصنوعی و آموزش علوم، فناوری، مهندسی و ریاضیات را گسترش داده اند، حال آن که سیاست دولتی کاربرد هوش مصنوعی را در رشته ها و صنایع گوناگون ترویج می کند. این به چین خط تولید بسیار بزرگی از مهندسان و پژوهشگران می دهد. صرفه جویی های مقیاس، کمبود متخصصان بسیار باتجربه را از میان نمی برد، و گسترش سریع نیز کیفیت را تضمین نمی کند. با این همه، راهبرد آموزشی بخشی از همان واکنش گسترده تر به محدودیت های فناورانه است. جایی که دسترسی به سخت افزار خط مقدم محدود است، چین می کوشد از طریق راه حل های مهندسی، بهره وری نرم افزاری، زیرساخت، آموزش و سرمایه ی انسانی جبران کند.

بنابراین کشمکش آمریکا-چین با رقابت معمولی میان شرکت های فردی تفاوت دارد. هر دو دولت به طور فزاینده کنترل بر نیمه هادی ها، مدل های هوش مصنوعی، مراکز داده، منابع انرژی و نیروی متخصص فنی را اجزای قدرت ملی می دانند.

محدودیت های صادراتی که برای کند کردن یک رقیب طراحی شده اند می توانند جانشینی داخلی را برانگیزند. راهبردهای وزن باز می توانند رقبای اختصاصی مستقر را به چالش کشند. پیشرفت های فناورانه ی تجاری می توانند اهمیت یی نظامی بیابند. هر طرف می تواند شتاب بیشتر خود را با این استدلال توجیه کند که طرف دیگر همان کار را می کند. این مسابقه به خودی خود تقویت می شود.

هوش مصنوعی نظامی

هرچه تا اینجا بحث شد - رقابت شرکتی، سرمایه گذاری دولتی، محدودیت های نیمه هادی و مسابقه برای ظرفیت محاسباتی - پیامدهای مستقیمی برای کاربردهای نظامی هوش مصنوعی دارد. هوش مصنوعی هم اکنون در تحلیل اطلاعاتی، نظارت،

خطر اصلی سرمایه‌داری است، نه هوش مصنوعی

عملیات سایبری، پشتیبانی لجستیک، پهپادهای خودمختار و نیمه‌خودمختار، جنگ الکترونیک و تصمیم‌گیری نظامی نقش دارد.

نسخه‌ی علمی-تخیلی خطر، ارتشی از ربات‌ها است که خودبه‌خود تصمیم به آغاز جنگ می‌گیرند. خطر واقعی‌تر و کم‌تر نمایشی این است: ماشین‌ها به تدریج در مراحل بیشتری از فرایند تشخیص یک اتفاق تا تصمیم درباره‌ی پاسخ به آن جای می‌گیرند. سیستم‌های هوش مصنوعی می‌توانند اطلاعات از سنسورهای گوناگون را ترکیب کنند، هدف‌های ممکن را شناسایی کنند، تهدیدها را اولویت‌بندی کنند، اقدام‌ها را توصیه کنند و شمار زیادی سکو را هماهنگ کنند. یکی از جذابیت‌های چنین سیستم‌هایی دقیقاً سرعت آن‌هاست. اما همین سرعت خطر خود را می‌آفریند. هرچه تصمیم‌گیری نظامی سریع‌تر شود، انسان‌ها ممکن است زمان کم‌تری برای زیر سؤال بردن اطلاعاتی که دریافت می‌کنند، بازاندیشی درباره‌ی یک تفسیر یا توقف زنجیره‌ای رو به تشدید از حوادث داشته باشند.

این مسئله‌ی آشنای مسابقه‌ی تسلیحاتی را پدید می‌آورد. تصور کنید آمریکا، چین و بریتانیا همگی می‌پذیرند که اجازه دادن به ماشین‌ها برای تصمیم‌های خودمختار مرگ‌وزندگی می‌تواند خطرناک باشد. این درک مشترک لزوماً هیچ‌یک از آن‌ها را از توسعه‌ی این فناوری باز نمی‌دارد. اگر هر دولت باور کند رقابیش تصمیم‌گیری نظامی را خودکار می‌کنند و از این راه برتری کسب می‌کنند، هر کدام محرکی برای توسعه‌ی سیستم‌های مشابه دارد صرفاً تا از قافله عقب نماند. هیچ‌کس لازم نیست پیامد خطرناک را بخواهد. فشار رقابتی می‌تواند آن را پدید آورد.

توانمندی‌های سایبری مبتنی بر هوش مصنوعی بُعد دیگری می‌افزایند. سیستم‌هایی که می‌توانند آسیب‌پذیری‌ها را شناسایی کنند، نرم‌افزار بنویسند و رشته‌ای از کنش‌ها را در سراسر شبکه‌های رایانه‌ای انجام دهند می‌توانند برخی عملیات سایبری را به مراتب سریع‌تر از تیم‌های انسانی انجام دهند. این شبکه‌ها محدود به نظامی نیستند. بیمارستان‌ها، شبکه‌های برق، مؤسسات مالی، سیستم‌های ارتباطی و دیگر زیرساخت‌های غیرنظامی به سیستم‌های رایانه‌ای به هم پیوسته وابسته‌اند.

همچنین هیچ مرز فناورانه‌ی روشنی میان «هوش مصنوعی غیرنظامی» و «هوش مصنوعی نظامی» وجود ندارد. همان نوع پردازنده‌های پیشرفته‌ای که برای آموزش یک

چتبات تجاری به کار می‌روند می‌توانند برای مدل‌های نظامی نیز به کار روند. فناوری تشخیص تصویری که می‌تواند محصولات معیوب را در خط تولید کارخانه شناسایی کند، می‌تواند به شناسایی اهداف نظامی نیز یاری رساند. یک معماری عامل هوش مصنوعی که برای انجام وظایف متداول یک شرکت طراحی شده می‌تواند برای عملیات سایبری نیز اقتباس شود.

بنابراین رقابت تجاری و نظامی یکدیگر را تغذیه می‌کنند. شرکت‌ها سیستم‌های همه‌کاره‌ی روبه‌رشدی می‌سازند. دولت‌ها و نیروهای نظامی کاربردهای راهبردی ممکن را شناسایی می‌کنند و بودجه را به سوی تطبیق این سیستم‌ها هدایت می‌کنند. تقاضای نظامی و اطلاعاتی سپس سرمایه‌گذاری و توسعه‌ی فنی بیشتری فراهم می‌کند که بخش تجاری می‌تواند از آن سود ببرد.

از این‌رو هوش مصنوعی به جای صرفاً یک فناوری مصرفی با کاربردهای نظامی فرعی، بخشی از نظام صنعتی-نظامی می‌شود. بحث درباره‌ی هوش مصنوعی اغلب به «بشریت» ارجاع می‌دهد که به آستانه‌ی فناوری‌های فناورانه‌ای نزدیک می‌شود، گویی بشریت یک کنش‌گر واحد با یک مجموعه‌ی منافع باشد. چنین نیست. شرکت‌های فناوری، کارکنانشان، سرمایه‌گذاران، دولت‌ها، نیروهای نظامی، سازمان‌های اطلاعاتی، دانشگاه‌ها و کاربران عادی، درباره‌ی اینکه کدام سیستم‌ها توسعه می‌یابند، چگونه به کار گرفته می‌شوند و به چه اهدافی خدمت می‌کنند، قدرت به‌شدت نابرابری دارند.

بنابراین پرسش سودمندتر صرفاً این نیست که «آیا بشریت می‌تواند هوش مصنوعی را کنترل کند؟» بلکه این است: «چه کسی هوش مصنوعی را کنترل می‌کند، به نفع چه کسی، و پاسخ‌گو به چه کسی؟» از این زاویه، داستان مرکزی یک آینده‌ی فرضی نیست که در آن ماشین‌ها علیه انسان به‌پامی‌خیزند؛ بلکه اکنونی است که در آن فناوری‌های تولیدی فوق‌العاده قدرتمند درون نظامی اقتصادی متشکل بر مالکیت خصوصی و رقابت توسعه می‌یابند، در حالی که دولت‌ها هم‌زمان برای برتری ژئوپلیتیکی و نظامی رقابت می‌کنند.

^۱ اشاره به پرتاب اسپوتنیک ۱ توسط اتحاد جماهیر شوروی در سال ۱۹۵۷، نخستین ماهواره‌ی مصنوعی جهان. اسپوتنیک ایالات متحده را غافلگیر کرد، زیرا نشان داد رقیبی فناورانه که به‌طور گسترده تصور می‌شد از قافله عقب است، به یک پیشرفت بزرگ دست یافته است.