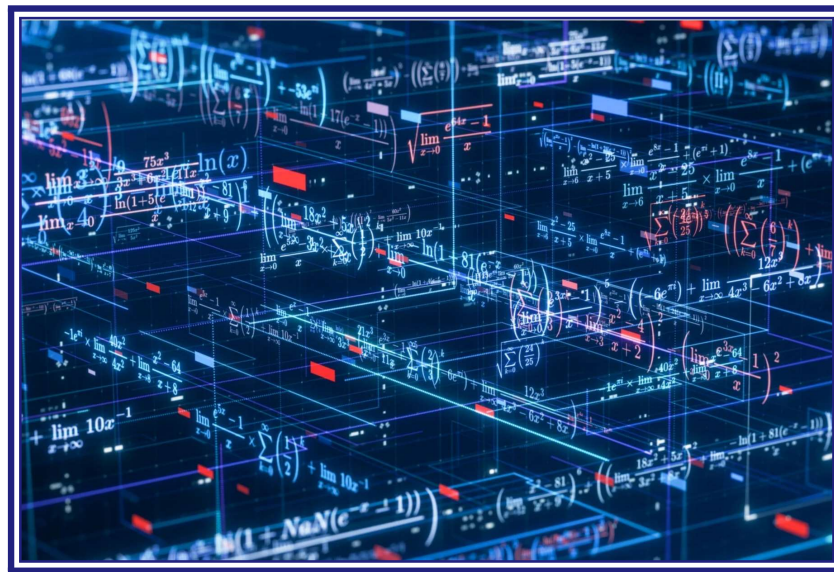




هوش مصنوعی مدال المپیاد ریاضی را به دست گرفت - اما ریاضی دانان تحت تأثیر قرار نگرفتند

امیلی ریل **

ترجمه: سمیه رهنما *



المپیاد امسال که ماه گذشته در ساحل سان‌شاین استرالیا برگزار شد، یک حاشیه‌ی غیرمعمول داشت. در حالی که ۱۱۰ دانش‌آموز از سراسر جهان با قلم و کاغذ مشغول حل مسائل پیچیده بودند، چند شرکت فعال در حوزه‌ی هوش مصنوعی، مدل‌های جدید خود را به صورت آزمایشی روی نسخه‌ی دیجیتالی آزمون امتحان کردند. بلافاصله پس از مراسم اختتامیه، شرکت‌های OpenAI و Google DeepMind اعلام کردند که مدل‌هایشان موفق به حل پنج مسئله از شش مسئله شده‌اند و مدال طلای غیررسمی کسب کرده‌اند. برخی پژوهشگران مانند سباستین بوبک^۱ از OpenAI این موفقیت را «لحظه‌ی فرود بر ماه» برای صنعت توصیف کردند.

آیا هوش مصنوعی جای ریاضی‌دانان را می‌گیرد؟

اما آیا واقعاً چنین است؟ آیا هوش مصنوعی قرار است جایگزین

مدل‌های هوش مصنوعی ظاهراً در حل مسائل المپیاد بین‌المللی ریاضی عملکرد خوبی داشتند، اما نحوه‌ی رسیدن آن‌ها به پاسخ‌ها یادآور این نکته است که هنوز به افرادی که ریاضی می‌خوانند، نیاز داریم.

خاطره‌ای از دوران دبیرستان

یکی از خاطرات ماندگار من از سال آخر دبیرستان، شرکت در آزمون نه‌ساعته‌ای بود که تنها شامل شش سؤال ریاضی بود. شش نفر از برترین شرکت‌کنندگان به تیم ایالات متحده برای المپیاد بین‌المللی ریاضی (IMO) راه یافتند؛ قدیمی‌ترین رقابت ریاضی برای دانش‌آموزان دبیرستانی در جهان. من نتوانستم به آنجا راه پیدا کنم، اما بعدها استاد تمام ریاضی شدم.

حضور بی‌سروصدای هوش مصنوعی در رقابت

^۱ Sébastien Bubeck

لندن، در یک انجمن آنلاین گفت: «وقتی به عنوان دانشجوی کارشناسی وارد کمبریج شدم و مدال طلای IMO را در دست داشتم، در موقعیتی نبودم که بتوانم به هیچ یک از ریاضی دانان پژوهشگر در آنجا کمک کنم»

تعامل با مدل های زبانی

امروزه پژوهش های ریاضی گاهی بیش از یک عمر زمان می برد تا تخصص لازم برای آن کسب شود. مانند بسیاری از همکارانم، وسوسه شده ام که «اثبات حسی»^۴ را امتحان کنم. یک گفتگوی ریاضی با یک مدل زبانی بزرگ (LLM) درست مانند گفتگو با یک همکار، با این سوال که «آیا این درست است که...؟» و مدل نیز استدلالی واضح ارائه دهد. تجربه ای من نشان داده که این استدلال ها در موضوعات استاندارد معمولاً درست هستند، اما در مرزهای دانش اغلب اشتباهات ظریفی دارند. برای مثال، همه ی مدل هایی که من امتحان کرده ام، اشتباهی مشابه در نظریه ی خودتوان ها در کتگوری بی نهایت بعدی ضعیف مرتکب شده اند؛ فرض می کند نظریه ی خودتوان ها در مورد کتگوری های بی نهایت بعدی ضعیف همان گونه رفتار می کند که در مورد کتگوری عادی^۵ رفتار می کند. چیزی که متخصصان انسانی در رشته من می دانند نادرست است.

اعتماد به اثبات های رسمی

من هرگز به یک مدل زبانی - که در اصل فقط پیش بینی کننده ی کلمات بعدی بر اساس داده هایش است - برای ارائه ی اثباتی که خودم نتوانم بررسی کنم، اعتماد نخواهم کرد. خبر خوب این است که ابزارهایی به نام «دستیار اثبات» وجود دارند؛ برنامه هایی که (بدون استفاده از هوش مصنوعی) بررسی می کنند آیا یک استدلال منطقی واقعاً ادعای مطرح شده را اثبات می کند یا نه. این ابزارها توجه ریاضی دانانی مانند تائو، بازارد و خود من را که خواهان اطمینان بیشتری از صحت اثبات های خود هستیم، جلب کرده اند و می توانند به دموکراتیزه کردن ریاضیات و حتی افزایش ایمنی هوش مصنوعی کمک کنند.

مثال تاریخی: رامانوجان

فرض کنید نامه ای با دست خطی ناآشنا از شهر اُرد^۶ (شهری در تاملیل نادو) در هند دریافت کنم که حاوی اثباتی ریاضی باشد. شاید ایده های آن درخشان باشند یا شاید بی معنا. باید ساعت ها وقت صرف کنم تا خطبه خط آن را بررسی و درست و غلط بودن آن را کنترل کنم. یک ریاضی دان انسانی، مانند من، فقط باید معنای اصطلاحات فنی در

ریاضی دانان حرفه ای شود؟ من هنوز منتظر اثبات این ادعا هستم. هیجان زندگی و تبلیغات درباره ی نتایج هوش مصنوعی امسال قابل درک است، زیرا المپیاد بسیار دشوار است. در سال آخر دبیرستان، من حتی حساب دیفرانسیل و جبر خطی را کنار گذاشتم تا روی مسائل سبک المپیادی تمرکز کنم که چالش بیشتری داشتند. مدل های پیشرفته ی در حال توسعه، عملکرد بسیار بهتری نسبت به مدل های تجاری فعلی موجود در بازار داشتند. در رقابت موازی برگزار شده توسط MathArena.ai، مدل هایی مانند Gemini 2.5 Pro، Grok، O3 High 4، و O4-Mini High و DeepSeek R1 حتی یک پاسخ کاملاً درست ارائه نکردند. این نشان می دهد که مدل های هوش مصنوعی در حال هوشمندتر شدن هستند و قابلیت استدلال آن ها به طور چشمگیری در حال پیشرفت است.

مقایسه ی ناعادلانه

با این حال، من هنوز نگران نیستم. این مدل ها تنها در یک آزمون نمره ی خوبی گرفتند - همان طور که بسیاری از دانش آموزان نیز این طور بودند - و مقایسه ی مستقیم چندان منصفانه نیست. مدل ها اغلب از استراتژی «بهترین از میان چند» استفاده می کنند؛ یعنی چند پاسخ تولید کرده و خودشان به آن ها نمره می دهند، سپس بهترین را انتخاب می کنند. این شبیه به حالتی است که چند دانش آموز به طور مستقل کار کنند و سپس دور هم جمع شوند تا بهترین راه حل را انتخاب کرده و ارسال کنند. اگر به شرکت کنندگان انسانی چنین امکانی داده شود، احتمالاً نمرات آن ها نیز بهتر خواهد شد.

نظر ریاضی دانان برجسته

ریاضی دانان برجسته ی دیگر نیز نسبت به این هیاهو محتاط بوده و هشدار می دهند. ترنس تائو، دارنده ی مدال طلای المپیاد بین المللی ریاضی IMO (و در حال حاضر استاد دانشگاه کالیفرنیا، لس آنجلس UCLA)، در شبکه ی Mastodon اشاره کرد که توانایی های هوش مصنوعی به روش آزمون بستگی دارد. رئیس IMO، گرگور دولینار^۲ نیز اظهار داشت که این سازمان نمی تواند روش های مورد استفاده توسط مدل های هوش مصنوعی؛ از جمله میزان محاسبات، دخالت انسانی، یا اینکه آیا نتایج قابل تکرارند را تأیید کند.

تفاوت با پژوهش های واقعی

به علاوه، به نظر من سؤالات المپیاد بین المللی ریاضی (IMO) قابل مقایسه با مسائل پژوهشی که ریاضی دانان حرفه ای سعی در پاسخ دادن به آن ها دارند نیستند؛ مسائلی که گاهی حل آن ها ۹ ساعت بلکه ۹ سال زمان می برد. کوین بازارد^۳، استاد ریاضی در کالج سلطنتی

^۲Gregor Dolinar

^۳Kevin Buzzard

^۴vibe proving

^۵ordinary

^۶Erode

در حالی که مدل‌های «استدلالی» با هدف تجزیه مسائل به بخش‌های کوچک‌تر و توضیح مرحله به مرحله «تفکر» خود هدایت می‌شوند، خروجی آنها به همان اندازه ممکن است استدلال‌هایی تولید کنند که منطقی به نظر برسند، اما در واقع نادرست باشد و با یک اثبات واقعی اشتباه گرفته شود. در مقابل، دستیار اثبات تنها زمانی یک اثبات را می‌پذیرد که کاملاً دقیق و منسجم باشد و هر مرحله از زنجیره فکری خود را توجیه کند. در برخی موارد، راه‌حل تقریبی کافی است، اما وقتی دقت ریاضی اهمیت دارد، باید از مدل‌های هوش مصنوعی بخواهیم که اثبات‌هایشان به طور رسمی قابل اعتبارسنجی باشند. در زندگی، عدم قطعیت فراوان است و اشتباه کردن آسان. همانطور که در دبیرستان آموختم، یکی از بهترین چیزها در مورد ریاضی این واقعیت است که می‌توانید به طور قطعی ثابت کنید که برخی ایده‌ها اشتباه هستند. بنابراین خوشحالم که یک هوش مصنوعی سعی می‌کند مسائل ریاضی شخصی‌ام را حل کند، اما تنها در صورتی که نتایج به طور رسمی قابل تأیید باشند و ما هنوز به آن نقطه نرسیده‌ایم.

*Emily Riehl, [AI Took on the Math Olympiad—But Mathematicians Aren't Impressed](#), Scientificamerican, August 7, 2025.

**دانشگاه یزد

عبارت قضیه را بفهمد. اما اگر متن ریاضی به زبان رایانه‌ای مناسب نوشته شده باشد، یک دستیار اثبات می‌تواند منطق آن را برایم بررسی کند. در مورد سرینواسا رامانوجان، نابغه‌ی ریاضی اهل اُرد، یک متخصص وقت گذاشت تا نامه‌اش را رمزگشایی کند. در سال ۱۹۱۳، رامانوجان به ریاضی‌دان بریتانیایی جی. ایچ. هاردی نامه نوشت و خوشبختانه هاردی نبوغ او را تشخیص داد و او را به کمبریج دعوت کرد و این آغاز دوران حرفه‌ای یکی از «بزرگان» ریاضی تمام دوران بود.

اثبات‌های رسمی و مدل‌های هوش مصنوعی

نکته‌ی جالب این است که برخی از مدل‌های هوش مصنوعی شرکت‌کننده در IMO پاسخ‌های خود را به زبان دستیار اثبات Lean ارائه دادند تا برنامه‌ی رایانه‌ای بتواند به طور خودکار خطاهای منطقی را بررسی کند. استارت‌آپی به نام Harmonic اثبات‌های رسمی برای پنج مسئله از شش مسئله منتشر کرد و ByteDance با حل چهار مسئله به سطح مدال نقره رسید. البته سؤالات باید متناسب با محدودیت‌های زبانی مدل‌ها طراحی می‌شدند و آنها هنوز به روزها زمان نیاز داشتند تا آن را حل کنند.

نتیجه‌گیری

با این حال اثبات‌های رسمی به طور منحصر به فرد قابل اعتماد هستند.

