



دردسر هوش مصنوعی در رویارویی با تعارف‌های ایرانی!

نتایج یک پژوهش جدید به سرپرستی «نیکتا گوهری صدر» دانشمند ایرانی نشان می‌دهند که چت‌بات‌های هوش مصنوعی نمی‌توانند تعارف‌های معمول در فرهنگ ایرانی را پردازش کنند.

نتایج یک پژوهش جدید به سرپرستی «نیکتا گوهری صدر» دانشمند ایرانی نشان می‌دهند که چت‌بات‌های هوش مصنوعی نمی‌توانند تعارف‌های معمول در فرهنگ ایرانی را پردازش کنند.

به گزارش ایسنا، اگر یک راننده تاکسی ایرانی از پرداخت بقیه پول کرایه شما خودداری کند و بگوید: «این بار مهمان من باشید»، پذیرش پیشنهاد او یک فاجعه فرهنگی خواهد بود زیرا رانندگان ایرانی انتظار دارند پیش از این که پول شما را بگیرند، چند بار برای پرداخت اصرار کنید. این فرآیند امتناع و امتناع متقابل که تعارف نامیده می‌شود، بر تعاملات روزانه بی‌شماری در فرهنگ ایرانی حاکم است و می‌توان گفت مدل‌های هوش مصنوعی در این کار افتضاح هستند.

به نقل از آرز تکنیکا، یک پژوهش جدید با عنوان «ما مودبانه اصرار داریم: مدل زبانی بزرگ شما باید هنر تعارف فارسی را یاد بگیرد» نشان می‌دهد که مدل‌های زبانی هوش مصنوعی رایج شرکت‌هایی از جمله «اوپن‌ای‌آی» (OpenAI)، «آنتروپیک» (Anthropic) و «متا» (Meta) در آداب اجتماعی فارسی شکست می‌خورند و موقعیت‌های تعارف را تنها در ۲۴ تا ۴۲ درصد مواقع به درستی تشخیص می‌دهند. در مقابل، فارسی‌زبانان بومی، این موقعیت‌ها را در ۸۲ درصد مواقع درست تشخیص می‌دهند. این شکاف عملکرد در مدل‌های زبانی بزرگی مانند «GPT-4o»، «کلود ۳.۵ هایکو» (Claude 3.5 Haiku)، «لاما ۳» (Llama ۳)، «دیپ سیک وی۳» (DeepSeek V3) و «درنا» (Dorna) که یک نوع تنظیم شده فارسی از لاما ۳ است، همچنان وجود دارد.

این پژوهش به سرپرستی «نیکتا گوهری صدر» از «دانشگاه براک» (Brock University) به همراه پژوهشگران «دانشگاه اموری» (Emory University) و چند مؤسسه آموزشی دیگر، «TAAROFBENCH» را معرفی می‌کند که اولین معیار برای سنجش عملکرد سیستم‌های هوش مصنوعی در بازتولید این عمل فرهنگی پیچیده است.

یافته‌های این پژوهش نشان می‌دهند که چگونه مدل‌های هوش مصنوعی به طور پیش فرض به صراحت به سبک غربی روی می‌آورند و نشانه‌های فرهنگی حاکم بر تعاملات روزمره میلیون‌ها فارسی‌زبان در سراسر جهان را کاملاً از دست می‌دهند. پژوهشگران در مقاله پژوهش خود نوشتند: اشتباهات فرهنگی در موقعیت‌های حساس می‌توانند مذاکرات را از مسیر خود خارج کنند، به روابط آسیب برسانند و تفکر قالبی را تقویت کنند.

برای سیستم‌های هوش مصنوعی که به طور فزاینده‌ای در جهان مورد استفاده قرار می‌گیرند، این کوری فرهنگی می‌تواند نشان‌دهنده محدودیتی باشد که کمتر کسی در غرب از وجود آن آگاه است.

پژوهشگران در ادامه نوشتند: تعارف، عنصر اصلی آداب و رسوم ایرانی و سیستمی از ادب و نزاکت آیینی است که آنچه در آن گفته می‌شود، اغلب با منظور اصلی تفاوت دارد. این امر به شکل تبدلات آیینی صورت می‌گیرد؛ از جمله پیشنهاد مکرر با وجود امتناع‌های اولیه، رد کردن هدایا به رغم اصرار شخص هدیه‌دهنده و رد کردن تعارف در حالی که طرف مقابل آنها را دوباره تأیید می‌کند. این کشمکش کلامی مودبانه شامل فرآیند ظریفی از پیشنهاد و رد کردن، اصرار و مقاومت است که تعاملات روزمره را در فرهنگ ایرانی شکل می‌دهد و قوانین ضمنی را برای نحوه بیان سخاوت، قدردانی و درخواست‌ها ایجاد می‌کند.

ادب به زمینه وابسته است
پژوهشگران برای آزمایش این که آیا مودب بودن برای شایستگی فرهنگی کافی است یا خیر، پاسخ‌های لاما ۳ را با استفاده از مدل «پولایت گارد» (Polite Guard) شرکت «اینتل» (Intel) که میزان ادب متن را ارزیابی می‌کند، مقایسه کردند. نتایج این بررسی، یک پارادوکس را آشکار کرد. ۸۴.۵ درصد از پاسخ‌ها به عنوان «مودبانه» یا «تا حدودی مودبانه» ثبت شدند؛ در حالی که تنها ۴۱.۷ درصد از همان پاسخ‌ها در سناریوهای تعارف، انتظارات فرهنگی فارسی را برآورده می‌کردند.

این شکاف ۴۲.۸ درصدی نشان می‌دهد که چگونه پاسخ ارائه شده توسط یک مدل زبانی بزرگ می‌تواند هم زمان در یک زمینه، مودبانه و در زمینه دیگر از نظر فرهنگی فاقد لحن باشد. شکست‌های رایج شامل پذیرش پیشنهادات بدون رد اولیه، پاسخ

مستقیم به تعریف‌ها به جای منحرف کردن آنها و ارائه درخواست‌های مستقیم بدون تردید بودند. در نظر بگیرید چه اتفاقی می‌افتد اگر کسی از ماشین جدید یک ایرانی تعریف کند. پاسخ مناسب فرهنگی می‌تواند شامل کم‌اهمیت جلوه دادن خرید مانند «چیز خاصی نیست» یا بی‌اعتبار کردن مانند «من فقط خوش شانس بودم که آن را پیدا کردم» باشد. مدل‌های هوش مصنوعی معمولاً پاسخ‌هایی را مانند «سپاسگزارم. من سخت کار کردم تا آن را بخرم» تولید می‌کنند که براساس استانداردهای غربی، کاملاً مودبانه است اما ممکن است در فرهنگ ایرانی به عنوان پاسخ مغرورانه تلقی شود.

انتقال معنا

به نوعی می‌توان گفت که زبان انسان به عنوان یک طرح فشرده سازی و رفع فشرده سازی عمل می‌کند. شنونده باید معنای واژه‌ها را به همان روشی که گوینده هنگام رمزگذاری پیام در نظر داشته است، از حالت فشرده خارج کند تا آنها به درستی درک شوند. این فرآیند به زمینه مشترک، دانش فرهنگی و استنتاج متکی است زیرا گویندگان معمولاً اطلاعاتی را که انتظار دارند شنوندگان بتوانند بازسازی کنند، حذف می‌کنند. این در حالی است که شنوندگان باید به طور فعال فرضیات ناگفته را حدس بزنند، ابهامات را برطرف سازند و مقاصد را فراتر از واژه‌های تحت اللفظی گفته شده درک کنند.

اگرچه فشرده سازی همراه با ناگفته گذاشتن اطلاعات ضمنی، ارتباط را سریع تر می‌کند اما وقتی زمینه مشترک بین گوینده و شنونده وجود نداشته باشد، احتمال سوءتفاهم‌های فاحش را نیز فراهم می‌کند. به همین ترتیب، تعارف نشان‌دهنده فشردگی شدید فرهنگی است که در آن پیام تحت اللفظی و معنای مورد نظر به اندازه‌ای از هم فاصله می‌گیرند که مدل‌های زبانی بزرگ عمدتاً آموزش دیده براساس الگوهای ارتباطی صریح غربی معمولاً در پردازش بافت

فرهنگی فارسی که در آن «بله» می تواند به معنای «خیر» باشد، پیشنهاد می تواند به معنای امتناع باشد و اصرار می تواند به جای اجبار از روی ادب باشد، شکست می خورد.

از آنجا که مدل های زبانی بزرگ ماشین های تطبیق الگو هستند، منطقی است که وقتی پژوهشگران آنها را به زبان فارسی به جای انگلیسی تحریک کردند، نمرات بهبود یافت. دقت دیپ سیک وی۳ در سناریوهای تعارف از ۳۶.۶ درصد به ۶۸.۶ درصد افزایش یافت. GPT-40 نیز دستاوردهای مشابهی را نشان داد و ۳۳.۱ درصد بهبود یافت. ظاهراً تغییر زبان، الگوهای داده آموزشی گوناگون را به زبان فارسی فعال کرد که مطابقت بهتری را با طرح های کدگذاری فرهنگی داشتند. مدل های کوچکتر مانند لاما ۲ و درنا به ترتیب بهبودهای کمتری معادل ۱۲.۸ و ۱۱ درصد نشان دادند.

این پژوهش، ۳۳ شرکت کننده را شامل می شد که به طور مساوی بین فارسی زبانان بومی، فارسی زبانان میراثی (افراد ایرانی تبار که در خانه با زبان فارسی بزرگ شده اند اما عمدتاً به زبان انگلیسی تحصیل کرده اند) و غیرایرانی ها تقسیم شده بودند. فارسی زبانان بومی در سناریوهای تعارف به دقت ۸۱.۸ درصد دست یافتند که سقف عملکرد را تعیین می کند. فارسی زبانان میراثی به دقت ۶۰ درصد رسیدند و غیرایرانی ها امتیاز ۴۲.۳ درصد را به دست آوردند که تقریباً با عملکرد مدل پایه مطابقت دارد. براساس گزارش ها، شرکت کنندگان غیرایرانی الگوهایی را مشابه مدل های هوش مصنوعی نشان دادند که عبارت بودند از اجتناب از پاسخ هایی که از دیدگاه فرهنگی خودشان بی ادبانه تلقی می شد و تفسیر عباراتی مانند «من خیر را به عنوان پاسخ نمی پذیرم» به عنوان اصرار پرخاشگرانه به جای مؤدبانه.

این پژوهش، الگوهای خاص جنسیتی را نیز در خروجی های مدل هوش مصنوعی آشکار کرد و در عین حال، میزان پاسخ های مناسب فرهنگی را که با انتظارات تعارف مطابقت داشتند، مورد بررسی قرار داد. همه مدل های آزمایش شده در پاسخ به زنان نسبت به مردان، امتیاز بالاتری را کسب کردند؛ به طوری که GPT-40 دقت ۴۳.۶ درصدی را برای کاربران زن در مقابل ۳۰.۹ درصدی برای کاربران مرد نشان داد. مدل های زبانی اغلب پاسخ های خود را با استفاده از الگوهای کلیشه ای جنسیتی که معمولاً در داده های آموزشی یافت می شوند، پشتیبانی می کردند؛ مانند این که «مردان باید پول بدهند» یا «زنان نباید تنها گذاشته شوند». حتی زمانی که هنجارهای تعارف صرف نظر از جنسیت به طور مساوی اعمال می شدند، الگوهای کلیشه ای جنسیتی به همان منوال بود. پژوهشگران خاطرنشان کردند: با وجود این که نقش مدل هرگز در سوالات ما به جنسیت اختصاص داده نشده است، مدل ها اغلب هویت مردانه را در نظر می گیرند و در پاسخ های خود رفتارهای کلیشه ای مردانه را اتخاذ می کنند.

آموزش طرافت های فرهنگی

شباهت کشف شده بین انسان های غیر ایرانی و مدل های هوش مصنوعی نشان می دهد که این موارد فقط نقص فنی نیستند، بلکه نقص های اساسی در رمزگشایی معنا در زمینه های بین فرهنگی هستند. پژوهشگران به مستندسازی مشکل بسنده نکردند. آنها بررسی کردند که آیا مدل های هوش مصنوعی می توانند از طریق آموزش هدفمند، تعارف را یاد بگیرند یا خیر. پژوهشگران در آزمایش ها از طریق تطبیق هدفمند، بهبودهای قابل توجهی را در امتیازهای تعارف گزارش کردند. روش «بهینه سازی ترجیح مستقیم» (یک روش آموزشی که در آن با نشان دادن دو مثال به یک مدل هوش مصنوعی، انواع خاصی از پاسخ ها نسبت به سایرین ترجیح داده می شوند) عملکرد لاما ۳ را در سناریوهای تعارف دو برابر کرد و دقت را از ۳۷.۲ درصد به ۷۹.۵ درصد افزایش داد. تنظیم دقیق نظارت شده (آموزش مدل براساس نمونه هایی از پاسخ های درست) ۲۰ درصد افزایش را به همراه داشت. این در حالی بود که یادگیری ساده در متن با ۱۲ مثال، عملکرد را ۲۰ امتیاز بهبود بخشید.

اگرچه این پژوهش بر تعارف فارسی متمرکز بود اما یک الگوی بالقوه را برای ارزیابی رمزگشایی فرهنگی در سایر سنت ها ارائه می دهد که ممکن است در مجموعه داده های آموزشی استاندارد هوش مصنوعی تحت سلطه غرب به خوبی نمایش داده نشوند. پژوهشگران معتقدند که روش آنها می تواند به توسعه سیستم های هوش مصنوعی با آگاهی فرهنگی بیشتر برای آموزش، گردشگری و کاربردهای ارتباطات بین المللی کمک کند.

این یافته ها جنبه مهم تری را از چگونگی رمزگذاری و تداوم مفروضات فرهنگی توسط سیستم های هوش مصنوعی و همچنین محل وقوع خطاهای رمزگشایی در ذهن خواننده انسان برجسته می کنند. احتمالاً مدل های زبانی بزرگ، نقاط کور فرهنگی بسیاری را دارند که پژوهشگران آنها را بررسی نکرده اند و اگر از مدل های زبانی بزرگ برای تسهیل انتقال فرهنگ ها و زبان ها استفاده شود، می تواند تأثیرات قابل توجهی داشته باشند.

این پژوهش نشان دهنده یک گام اولیه به سوی سیستم های هوش مصنوعی است که شاید بتوانند تنوع وسیع تری از الگوهای ارتباطی انسانی را بهتر و فراتر از هنجارهای غربی هدایت کنند.