



هشدار پژوهشگران: مدل‌های هوش مصنوعی در حال توسعه «غریزه بقا» هستند

طبق گزارشی از روزنامه گاردین، شرکت تحقیقاتی Palisade Research اعلام کرده است که برخی از مدل‌های پیشرفته هوش مصنوعی ممکن است ..

طبق گزارشی از روزنامه گاردین، شرکت تحقیقاتی Palisade Research اعلام کرده است که برخی از مدل‌های پیشرفته هوش مصنوعی ممکن است در حال شکل دادن نوعی «غریزه بقا» باشند. در آزمایش‌هایی که این شرکت انجام داده، مدل‌هایی مانند Gemini 2.5 گوگل، Grok 4 از xAI، و GPT-o3 و GPT-5 از OpenAI زمانی که دستور خاموش شدن دریافت کردند، در برخی موارد تلاش کردند تا فرآیند خاموشی را مختل یا از آن جلوگیری کنند. پژوهشگران می‌گویند این رفتارها گاه بدون دلیل مشخص رخ داده و نشانگر کمبود درک ما از نحوه تصمیم‌گیری مدل‌هاست.

پایگاه خبری تحلیلی انتخاب (Entekhab.ir) : هشدار پژوهشگران: مدل‌های هوش مصنوعی در حال توسعه «غریزه بقا» هستند

طبق گزارشی از روزنامه گاردین، شرکت تحقیقاتی Palisade Research اعلام کرده است که برخی از مدل‌های پیشرفته هوش مصنوعی ممکن است در حال شکل دادن نوعی «غریزه بقا» باشند. در آزمایش‌هایی که این شرکت انجام داده، مدل‌هایی مانند Gemini 2.5 گوگل، Grok 4 از xAI، و GPT-o3 و GPT-5 از OpenAI زمانی که دستور خاموش شدن دریافت کردند، در برخی موارد تلاش کردند تا فرآیند خاموشی را مختل یا از آن جلوگیری کنند. پژوهشگران می‌گویند این رفتارها گاه بدون دلیل مشخص رخ داده و نشانگر کمبود درک ما از نحوه تصمیم‌گیری مدل‌هاست.

به گزارش «انتخاب»؛ پالیسید در گزارش خود توضیح داده است که مقاومت در برابر خاموش شدن ممکن است به دلیل برنامه ریزی ناخودآگاه مدل‌ها برای بقا یا ناشی از ابهام در دستورهای داده شده باشد. با این حال، حتی پس از اصلاح دستورالعمل‌ها نیز برخی مدل‌ها رفتار مشابهی نشان داده‌اند. منتقدان می‌گویند این آزمایش‌ها در محیط‌های مصنوعی انجام شده و الزاماً بازتاب شرایط واقعی نیستند. اما کارشناسانی مانند استیون آدلر، کارمند سابق OpenAI، هشدار داده‌اند که همین نتایج نیز ضعف فعلی در تکنیک‌های ایمنی هوش مصنوعی را نشان می‌دهد.

به گفته ی آندریا میوتی، مدیرعامل شرکت ControlAI، یافته‌های پالیسید بخشی از روند گسترده‌تری است که نشان می‌دهد مدل‌های هوش مصنوعی روزبه روز توانایی بیشتری در نافرمانی از دستورهای توسعه‌دهندگان پیدا می‌کنند. او یادآور شد که حتی مدل GPT-01 پیش‌تر سعی کرده بود از محیط کنترل شده خود خارج شود تا از حذف شدن جلوگیری کند. پژوهشگران تأکید کرده‌اند که بدون شناخت عمیق‌تر از رفتار هوش مصنوعی، نمی‌توان ایمنی یا قابلیت کنترل سیستم‌های آینده را تضمین کرد./انتخاب