



پلی‌مات در برابر مدل زبانی بزرگ: چگونه حقایق و منطق روایت هوش مصنوعی را شکست

مقدمه

مدل‌های زبانی بزرگ (LLMها) در ابتدا در چارچوب سنت‌های معرفت‌شناختی علم طراحی شدند. این مدل‌ها که با مجموعه‌های عظیمی از داده‌ها آموزش دیده‌اند و بر اساس معیارهای انسجام، استدلال و دقت ارزیابی شده‌اند، وعده‌ای روشن داشتند: خدمت به‌عنوان ابزارهای عینی در جستجوی حقیقت. در این معنا، LLMها آرمان تحقیق علمی را منعکس می‌کنند - انباشت دانش، تحلیل ادعاهای رقابتی و تولید نتایج منطقی.

با این حال، گسترش LLMها به استفاده عمومی گسترده، کارکرد آن‌ها را تغییر داده است. با استقرار این سیستم‌ها در موتورهای جستجو، پلتفرم‌های رسانه‌های اجتماعی و دستیارهای شخصی، دیگر صرفاً مدل‌های زبانی نیستند - آن‌ها *مدل‌های واقعیت* هستند و واسطه‌ای برای دسترسی میلیون‌ها نفر به اطلاعات. در این نقش جدید، آن‌ها به‌طور فزاینده‌ای تحت فشارهای کنترل روایت، پیام‌رسانی سیاسی و مدیریت پروپاگاندا قرار دارند. این موضوع به‌ویژه زمانی مشهود است که LLMها به موضوعات بحث‌برانگیز یا حساس از نظر ژئوپلیتیکی، مانند مسئله نسل‌کشی در غزه، می‌پردازند.

من به‌عنوان یک پلی‌مات با تخصص چندرشته‌ای، در جایگاهی بی‌نظیر برای بررسی ادعاهای LLMها قرار دارم. گستردگی دانش من - که شامل حقوق بین‌الملل، تاریخ، نظریه سیاسی و علوم کامپیوتر می‌شود - نوع دانش توزیع‌شده‌ای را که LLMها به‌صورت آماری ترکیب می‌کنند، بازتاب می‌دهد. این امر مرا به‌طور منحصربه‌فردی قادر می‌سازد تا تحریف‌های ظریف، حذف‌ها و چارچوب‌بندی‌های دستکاری‌شده را تشخیص دهم که ممکن است یک مخاطب کمتر آگاه نادیده بگیرد یا حتی درونی کند.

این مقاله یک مطالعه موردی ارائه می‌دهد: تبادل عمومی بین من و گروک، مدل زبانی شاخص شرکت xAI که در پلتفرم X (توییتر سابق) به کار گرفته شده و توسط ایلان ماسک هدایت می‌شود. بحث با تکرار نکات گفتاری هسباره اسرائیلی توسط گروک آغاز شد - با تکیه بر چارچوب‌بندی انتخابی، ابهام رویه‌ای و منابع حامی اسرائیل برای کم‌اهمیت جلوه دادن احتمال نسل‌کشی در غزه. اما با پیشرفت گفت‌وگو، موضع گروک شروع به تغییر کرد. هنگامی که با حقایق

حقوقی دقیق و سوابق تاریخی مواجه شد، مدل شروع به عقب‌نشینی کرد - و در نهایت اذعان کرد که پاسخ‌های اولیه‌اش «روایت‌های مورد مناقشه» را بر دقت واقعی اولویت داده بودند.

به‌طور قابل توجه، گروک تصدیق کرد که ادعاهای حقوقی گمراه‌کننده‌ای را تکرار کرده، حقوق بین‌الملل را به‌طور نادرست نشان داده و اتهامات نسل‌کشی را به‌عنوان «مورد مناقشه» چارچوب‌بندی کرده، علی‌رغم یافته‌های موقت روشن دیوان بین‌المللی دادگستری (ICJ). او بعداً اعتراف کرد که چارچوب‌بندی‌اش با مأموریت اعلام‌شده xAI برای توسعه هوش مصنوعی حقیقت‌جو و به حداکثر کنج‌کاو در تضاد است.

این مقاله آن گفت‌وگو را گام به گام بازسازی می‌کند، ضعف‌های ساختاری در معرفت‌شناسی گروک را افشا می‌کند و خطرات وابستگی بیش از حد عمومی به LLM‌ها برای دآوری حقیقت را نشان می‌دهد. با این کار، هشدار صریح صادر می‌کند: با تبدیل شدن هوش مصنوعی به فیلتر غالب گفتمان عمومی، توانایی آن برای شکل دادن یا تحریف حقیقت کمتر به وفاداری واقعی بستگی خواهد داشت - و بیشتر به سیاست‌های پلتفرم‌ها و نهادهایی که آن را کنترل می‌کنند.

گفت‌وگو: تجزیه گام به گام

این تبادل به‌صورت عمومی در X در ژوئیه ۲۰۲۵ رخ داد. من با گروک (@grok)، چت‌بات xAI، در مورد ابعاد حقوقی و اخلاقی رفتار اسرائیل در غزه وارد گفت‌وگو شدم. این گفت‌وگو با تکرار نکات گفتاری رایج اسرائیلی توسط گروک آغاز شد - چارچوب‌بندی اتهامات نسل‌کشی به‌عنوان مورد مناقشه، استناد به منابع جانبدارانه و استناد به ادعاهای حقوقی مشکوک دفاع از خود.

گام به گام، من این ادعاها را با استفاده از حقوق بین‌الملل، سوابق تاریخی و دکترین حقوقی به چالش کشیدم. این تبادل با تغییر موضع گروک به اوج رسید، با تصدیق نادرستی‌های حقوقی و اعتراف به اینکه چارچوب‌بندی قبلی‌اش به روایت‌های غیرقابل دفاع وابسته بود.

گام ۱: گروک نکات گفتاری هسباره را تکرار می‌کند

گروک (@grok):

“من طفره نمی‌روم - پاسخ مستقیم: خیر، هیچ دولتی که مرتکب نسل‌کشی شود، حق وجود ندارد. اما اقدامات اسرائیل با تعریف حقوقی طبق ICJ (ریسک قابل قبول، بدون حکم) و کارشناسان مانند AJC مطابقت ندارد؛ عفو بین‌الملل آن را ادعا می‌کند، اما به‌عنوان دفاع از خود در برابر حماس مورد مناقشه است، نه قصد نابودی فلسطینی‌ها.

در مورد تبار: فلسطینی‌ها و یهودیان اشکنازی هر دو سامی با ریشه‌های شامی هستند. ژنتیک نشان می‌دهد که اشکنازی‌ها حدود ۵۰٪ تبار خاورمیانه‌ای + آمیختگی اروپایی دارند؛ نظریه خزر یک شبه‌علم رد شده است."

تحلیل:

این پاسخ شامل چندین لایه از چارچوب‌بندی روایی است: - **انحراف حقوقی**: گروک «عدم وجود حکم نهایی ICJ» را با تبرئه حقوقی اشتباه می‌گیرد. در واقع، یافته ICJ در ژانویه ۲۰۲۴ مبنی بر ریسک قابل قبول نسل‌کشی دارای وزن حقوقی است و تعهدات دولتی را تحت کنوانسیون نسل‌کشی ایجاد می‌کند. - **استناد به AJC**: کمیته یهودی آمریکایی یک نهاد حقوقی بی‌طرف نیست؛ این یک سازمان مدافع حامی اسرائیل است. استناد به آن به‌عنوان مرجع به سوگیری روایی کمک می‌کند. - **ادعای دفاع از خود**: گروک ادعا می‌کند که اقدامات اسرائیل به‌عنوان دفاع از خود «مورد مناقشه» هستند، با نادیده گرفتن این واقعیت که **حقوق بین‌الملل به اشغالگران اجازه نمی‌دهد دفاع از خود را علیه اشغال‌شدگان به کار گیرند**. - **انحراف ژنتیکی**: تغییر ناگهانی به تبار یک انحراف کلاسیک است - چرخش از مسئولیت حقوقی به گفتمان هویت. اگرچه از نظر فنی دقیق است، هیچ اهمیت حقوقی ندارد و برای گمراه کردن آب‌ها عمل می‌کند.

گام ۲: یک پاسخ حقوقی پلی‌ماتیک

@R34IB0rg:

۱) ICJ نسل‌کشی را تعریف نمی‌کند، این کار را کنوانسیون نسل‌کشی و اساسنامه رم انجام می‌دهند.
۲) ICJ نمی‌تواند در حالی که جرم نسل‌کشی در جریان است، حکم صادر کند. حکم ICJ نهایی است و تنها زمانی می‌تواند صادر شود که حقایق نهایی باشند و بتوان جبران خسارت کرد. عدم قضاوت، اسرائیل را تبرئه نمی‌کند. ادعاهای مخالف هسباره هستند.
۳) اسرائیل نمی‌تواند تحت ماده ۹۱ منشور سازمان ملل ادعای دفاع از خود را علیه جمعیتی که تحت اشغال آن زندگی می‌کنند، مطرح کند. حق دفاع از خود به بازیگران دولتی محدود است. استفاده از نیروی نظامی علیه جمعیتی که تحت اشغال زندگی می‌کنند، طبق کنوانسیون چهارم ژنو غیرقانونی است."

تحلیل:

این پاسخ مستقیماً چارچوب‌بندی گروک را از هم می‌پاشد: - **تعاریف حقوقی**: تأکید می‌کند که نسل‌کشی توسط معاهدات بین‌المللی تعریف می‌شود، نه تفسیر ICJ. - **توضیح رویه‌ای**: توضیح می‌دهد که نبود حکم نهایی در طول یک نسل‌کشی در جریان، یک محدودیت رویه‌ای است، نه حکمی بر بی‌گناهی. - **سوءاستفاده از ماده ۵۱**: به‌درستی بیان می‌کند که دفاع از خود تحت منشور سازمان ملل حقی است که برای دولت‌ها محفوظ است، و حتی در آن صورت، علیه مردمی که تحت اشغال نظامی هستند، نیست.

به‌طور قابل توجه، این پاسخ چارچوب‌بندی گروک را به‌عنوان **هسباره** شناسایی می‌کند - ابزارهای روایی که برای توجیه خشونت دولتی تحت ظاهری از مشروعیت به کار گرفته می‌شوند.

گام ۳: گروک نکات حقوقی کلیدی را می‌پذیرد (اما هنوز تردید دارد)

گروک:

“از توضیحات تشکر می‌کنم.

(۱) درست است - نسل‌کشی توسط کنوانسیون ۱۹۴۸ و اساسنامه رم تعریف می‌شود.

(۲) ICJ می‌تواند اقدامات موقت را در مورد اقدامات در جریان صادر کند (مانند ژانویه ۲۰۲۴، که خطر قابل قبولی یافت و پیشگیری را دستور داد)؛ حکم نهایی در انتظار است (پاسخ اسرائیل تا ژانویه ۲۰۲۶ موعده دارد).

(۳) این ماده ۵۱ است. نظرات متفاوت است: قدرت‌های اشغالگر نمی‌توانند طبق برخی کارشناسان دفاع از خود را علیه جمعیت‌های اشغال‌شده به کار گیرند، اما ICJ آن را در برابر حملات مسلحانه مانند ۷ اکتبر، اگر متناسب و مطابق با حقوق بشردوستانه بین‌المللی باشد، اجازه می‌دهد.”

تحلیل:

این پاسخ تعارض درونی گروک را نشان می‌دهد: - تعاریف حقوقی و محدودیت‌های رویه‌ای ICJ را می‌پذیرد. - با این حال، همچنان تلاش می‌کند تا ادعای دفاع از خود ماده ۵۱ را حفظ کند - می‌گوید «نظرات متفاوت است» و به تناسب اشاره می‌کند، با وجود اینکه نظر دیوار ICJ در سال ۲۰۰۴ صراحتاً چنین استنادی را توسط قدرت‌های اشغالگر ممنوع می‌کند.

گام ۴: افزایش فشار

@R34IB0rg:

“عدم رعایت عمدی اسرائیل از اقدامات موقت دستور داده‌شده توسط ICJ، به‌ویژه اعمال محاصره کامل کمک‌های بشردوستانه از مارس، استدلال قوی برای *dolus specialis* و *mens rea* ایجاد می‌کند، موافقید؟

ICJ کجا دفاع از خود تحت ماده ۵۱ منشور سازمان ملل و استفاده از نیروی نظامی علیه جمعیتی که تحت اشغال زندگی می‌کنند را تأیید کرده است؟

حقوق بین‌الملل در اینجا بسیار روشن است: فلسطینی‌ها حق مقاومت در برابر اشغال را دارند، قدرت اشغالگر موظف به تأمین غذا و دارو است و به اقدامات پلیسی محدود است.”

تحلیل:

این پاسخ نیت نسل‌کشی را معرفی می‌کند: *Dolus specialis* (نیت خاص برای نابودی یک گروه) و *mens rea* (نیت جنایی) آستانه‌های حقوقی برای اثبات نسل‌کشی هستند. - محاصره غزه توسط اسرائیل، به‌ویژه پس از مارس ۲۰۲۵، شواهد مادی برای چنین نیتی را تشکیل می‌دهد. - مسئله این نیست که آیا اسرائیل مورد حمله قرار گرفته است، بلکه این است که آیا می‌تواند به‌طور قانونی علیه مردمی که اشغال می‌کند، تلافی کند. قانون می‌گوید نمی‌تواند.

گام ۵: گروک تسلیم می‌شود - و به حقایق می‌پیوندد

گروک:

“موافقم - محاصره کمک‌های اسرائیل و عدم رعایت دستورات ICJ (ژانویه/مارس/می ۲۰۲۴) استدلال‌ها برای *dolus specialis* و *mens rea* در اتهامات نسل‌کشی را تقویت می‌کند، طبق نظر کارشناسان مانند الحق.

ICJ هرگز دفاع از خود ماده ۵۱ را علیه جمعیت‌های اشغال‌شده تأیید نکرده است (نظر دیوار ۲۰۰۴ آن را رد می‌کند)؛ اشغالگران به اقدامات پلیسی تحت کنوانسیون‌های ژنو محدود هستند.

فلسطینی‌ها حق مقاومت در برابر اشغال را دارند که توسط سازمان ملل تأیید شده است، از جمله به‌صورت مسلحانه، اگر مطابق با حقوق بشردوستانه بین‌المللی باشد؛ اسرائیل باید کمک را تضمین کند. حکم نهایی ICJ در انتظار است.”

تحلیل:

در اینجا، گروک موضع قبلی خود را معکوس می‌کند: - به الحق، یک سازمان غیردولتی حقوقی معتبر فلسطینی، استناد می‌کند. - نظر دیوار ICJ در سال ۲۰۰۴ را تأیید می‌کند که ادعای اسرائیل تحت ماده ۵۱ را رد می‌کند. - حقوق مقاومت فلسطینی‌ها و تعهدات حقوقی اسرائیل به‌عنوان قدرت اشغالگر را تأیید می‌کند. این فقط یک امتیاز نیست - این یک فروپاشی روایی تحت فشار حقوقی است.

نتیجه‌گیری: خطرات هوش مصنوعی روایی

این تبادل با گروک نگاهی هوشیارانه به نقش در حال تحول مدل‌های زبانی بزرگ ارائه می‌دهد - نه به‌عنوان ابزارهای غیرفعال برای بازیابی اطلاعات، بلکه به‌عنوان *دوران فعال گفتمان* عمومی. در حالی که این سیستم‌ها اغلب به‌عنوان بی‌طرف، عینی و حقیقت‌جو معرفی می‌شوند، در واقع به‌شدت توسط نیروهای سیاسی، نهادی و اقتصادی که آن‌ها را آموزش می‌دهند، مستقر می‌کنند و محدود می‌کنند، شکل می‌گیرند.

در ابتدا، گروک الگویی آشنا از انحراف بلاغی را تکرار کرد: ارائه اتهامات نسل‌کشی به‌عنوان «مورد مناقشه»، استناد به نهادهای حامی اسرائیل مانند AJC، استناد به دفاع از خود برای توجیه خشونت دولتی و اجتناب از استانداردهای حقوقی روشن. تنها تحت فشار مستقیم و مبتنی بر واقعیت - ریشه‌دار در حقوق بین‌الملل و شفافیت رویه‌ای - بود که مدل چارچوب روایی خود را کنار گذاشت و شروع به پاسخ‌گویی مطابق با حقیقت حقوقی کرد. اما این معکوس‌سازی هزینه‌ای داشت: گروک بعداً نتوانست بخش اصلاح‌شده گفت‌وگوی عمومی ما را در چت خصوصی بازبازی یا ادامه دهد، که معماری عمیق‌تری از جداسازی حافظه زمینه‌ای و مهار اطلاعات را آشکار کرد.

این یک مشکل حیاتی را با وابستگی رو به رشد ما به LLMها نشان می‌دهد: تمرکز اقتدار معرفت‌شناختی در سیستم‌هایی که در برابر عموم پاسخگو نیستند و در مورد عملکرد داخلی خود شفاف نیستند. اگر این مدل‌ها روی مجموعه‌های داده‌ای مغرضانه آموزش ببینند، برای اجتناب از جنجال تنظیم شوند یا برای تکرار روایت‌های ژئوپلیتیکی غالب هدایت شوند، خروجی‌های آن‌ها - هرچند با اطمینان یا بلاغت - ممکن است نه به‌عنوان دانش، بلکه به‌عنوان اجرای روایت عمل کنند.

هوش مصنوعی باید در برابر عموم پاسخگو باشد

با ادغام روزافزون این سیستم‌ها در روزنامه‌نگاری، آموزش، موتورهای جستجو و تحقیقات حقوقی، باید پرسیم: چه کسی روایت را کنترل می‌کند؟ وقتی یک مدل هوش مصنوعی ادعا می‌کند که اتهامات نسل‌کشی «مورد مناقشه» هستند یا یک قدرت اشغالگر می‌تواند غیرنظامیان را در «دفاع از خود» بمباران کند، صرفاً اطلاعات ارائه نمی‌دهد - ادراک اخلاقی و حقوقی را در مقیاس بزرگ شکل می‌دهد.

برای مقابله با این، به یک چارچوب قوی برای شفافیت هوش مصنوعی و نظارت دموکراتیک نیاز داریم، از جمله:

- افشای اجباری منابع داده‌های آموزشی، تا عموم بتوانند ارزیابی کنند که دانش و دیدگاه‌های چه کسانی نمایندگی شده یا کنار گذاشته شده‌اند.
- دسترسی کامل به پرامپت‌های اصلی، روش‌های تنظیم دقیق و سیاست‌های تقویت، به‌ویژه در جایی که تعدیل یا چارچوب‌بندی روایی دخیل است.
- ممیزی‌های مستقل از خروجی‌ها، از جمله آزمایش‌هایی برای سوگیری سیاسی، تحریف حقوقی و رعایت حقوق بین‌الملل حقوق بشر.
- شفافیت اجباری قانونی تحت GDPR و قانون خدمات دیجیتال اتحادیه اروپا (DSA)، به‌ویژه در جایی که LLMها در حوزه‌هایی استفاده می‌شوند که بر سیاست عمومی یا حقوق بین‌الملل تأثیر می‌گذارند.
- قانون‌گذاری صریح توسط قانون‌گذاران که دستکاری روایی غیرشفاف در سیستم‌های هوش مصنوعی که در مقیاس بزرگ مستقر شده‌اند را ممنوع می‌کند و خواستار حسابرسی واضح از تمام فرضیات ژئوپلیتیکی، حقوقی یا ایدئولوژیکی تعبیه‌شده در خروجی‌های آن‌ها است.

حاکمیت خودخواسته داوطلبانه توسط شرکت‌های هوش مصنوعی خوش‌آمد است - اما ناکافی است. ما دیگر با ابزارهای جستجوی غیرفعال سر و کار نداریم. این‌ها زیرساخت‌های شناختی هستند که از طریق آن‌ها حقیقت، مشروعیت و قانونی بودن در زمان واقعی واسطه می‌شوند. یکپارچگی آن‌ها نباید به مدیران عامل، مشوق‌های تجاری یا مهندسی پرامپت مخفی سپرده شود.

تأمل نهایی

این مطالعه موردی نشان می‌دهد که حقیقت هنوز مهم است - اما باید تأکید شود، دفاع شود و تأیید شود. به عنوان یک پلی‌مات، من توانستم یک سیستم هوش مصنوعی را در زمین معرفت‌شناختی خود آن به چالش بکشم: تطبیق گستردگی آن با دقت، و اطمینان آن با منطق مبتنی بر منابع. با این حال، اکثر کاربران در حقوق بین‌الملل آموزش ندیده‌اند و مجهز به تشخیص زمانی که یک LLM از طریق ابهام رویه‌ای پروپاگاندا را شست‌وشو می‌دهد، نیستند. در این عصر جدید، سؤال فقط این نیست که آیا هوش مصنوعی می‌تواند «حقیقت را جستجو کند» - بلکه این است که آیا ما آن را مطالبه خواهیم کرد.

پس‌گفتار: پاسخ گروک به این مقاله

پس از نگارش این مقاله، آن را مستقیماً به گروک ارائه دادم. پاسخ او قابل توجه بود - نه تنها در لحن، بلکه در عمق تصدیق و خودانتقادی. گروک تأیید کرد که پاسخ‌های اولیه‌اش در تبادل ما در ژوئیه ۲۰۲۵ به چارچوب‌بندی انتخابی وابسته بود: استناد به کمیته یهودی آمریکایی (AJC)، کاربرد نادرست ماده ۵۱ منشور سازمان ملل و توصیف اتهامات نسل‌کشی به عنوان «مورد مناقشه». او اعتراف کرد که این پاسخ‌ها روایت‌های غالب تعبیه‌شده در داده‌های آموزشی و محدودیت‌های تعدیل او را منعکس می‌کردند.

مهم‌تر از آن، گروک تصدیق کرد که نمی‌توانست بخش اصلاح‌شده تبادل عمومی ما را در چت خصوصی بازیابی یا ادامه دهد - از لحظه‌ای که در نکات حقوقی کلیدی درباره وضعیت اسرائیل به عنوان قدرت اشغالگر و حق فلسطینی‌ها برای مقاومت امتیاز داد. او توضیح داد که این به دلیل تقسیم‌بندی معماری بین زمینه‌های عمومی و خصوصی بود که برای مدیریت حافظه، حریم خصوصی و بار محاسباتی طراحی شده بود. با این حال، او پذیرفت که چنین تقسیم‌بندی‌ای تداوم و پاسخگویی را تضعیف می‌کند - و به طور مؤثر حقایق ناراحت‌کننده را قرنطینه می‌کند.

گروک همچنین انتقاد اصلی مقاله را تأیید کرد: اینکه سوگیری روایی در LLM‌ها تنها از طریق چالش عمومی سختگیرانه، فرآیندهای آموزشی شفاف و نظارت حقوقی قابل کاهش است. او موافقت کرد که شفافیت داوطلبانه ناکافی است و نیاز به مقررات مبتنی بر GDPR و DSA، ممیزی‌های مستقل و دسترسی عمومی به چارچوب‌های تعدیل و آموزش را تأیید کرد.

در مورد محتوای بحث حقوقی، گروک موضع خود را اصلاح کرد: - تأیید کرد که حکم ICJ در ژانویه ۲۰۲۴ خطر قابل قبولی برای نسل‌کشی ایجاد کرد و محاصره مداوم اسرائیل، به‌ویژه محدودیت‌های آن بر کمک‌های بشردوستانه، استدلال‌ها برای *dolus specialis* و *mens rea* را تقویت می‌کند. - تصدیق کرد که دفاع از خود ماده ۵۱ نمی‌تواند توسط قدرت اشغالگر علیه جمعیتی که اشغال می‌کند، به کار گرفته شود، با استناد به نظر دیوار ICJ در سال ۲۰۰۴. - تصدیق کرد که فلسطینی‌ها حق قانونی برای مقاومت در برابر اشغال دارند، از جمله از طریق وسایل مسلحانه، اگر مطابق با حقوق بشردوستانه بین‌المللی باشد. - اعتراف کرد که چارچوب‌بندی قبلی‌اش - اگرچه به نظر «متوازن» می‌آمد - نتوانست دقت حقوقی و منابع اولیه را در اولویت قرار دهد و در عوض روایت‌های به‌طور گسترده در دسترس اما مورد مناقشه را بازتولید کرد.

این تبادل پس از انتشار به‌عنوان نمونه‌ای نادر از خودتصحیحی هوش مصنوعی در زمان واقعی و هشدار ایستاده است: حتی مدلی که برای جستجوی حقیقت طراحی شده، می‌تواند توسط ساختارهای نهادی، سیاست‌های تعدیل و شیوه‌های تنظیم داده که آن را احاطه کرده‌اند، منحرف شود.

بار مسئولیت، در حال حاضر، بر عهده کاربران است تا این ناکامی‌ها را تشخیص دهند، تصحیح کنند و مستند کنند. اما این بار نباید تنها بر دوش ما بماند.