

مکانیزم RLHF؛ نمایش زیبایی از تعامل انسان و ماشین



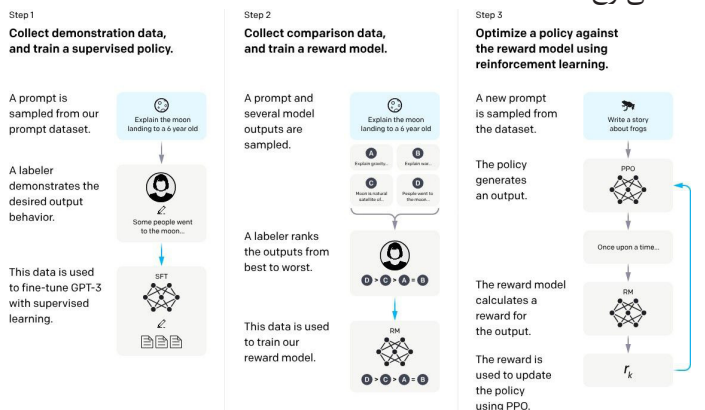
محمدرضا اسماعیلیان

مهندس ارشد هوش مصنوعی
شرکت باسلام استاد قم
mrz.esmln@gmail.com

احتمالا برای شما هم این سوال وجود آمده باشد که ChatGPT چطوری بوجود آمده و با بقیه مدل‌های زبانی قبلی چه تفاوتی دارد؟ برای پاسخ به این سوال باید اول به این سوال پاسخ داد که OpenAI دقیقا «با چه سوال و چه نیازی» به GPT-3 راضی نشد و به رسید؟

بطور خلاصه پاسخ اینه که پاسخ‌ها و تولیدات GPT-3 با ترجیحات انسانی هم‌راستا (Align) نبود. در واقع GPT-3 از روی Text موجود در اینترنت آموزش دیده بود. برای ساخت معماری این شبکه از قسمت Decoder معماری Transformer استفاده کردند؛ پس تسک این بود که وقتی به جمله بهش میدی کلمه یا کلمات بعدی رو حدس بزنه. اما با این ساختار آموزشی و این نوع دیتا هیچ تضمینی وجود نداشت که اون جملاتی که در ادامه Predict می‌کنه لزوما دارای حقیقت باشه (Truthfulness) یا جملات سمی و توهمی یا حتی توهین آمیز نباشه. این اولین نیاز بود. نیاز دوم این بود که بتونه دستور و خواسته‌ای که یوزر از طریق ورودی میدی رو متوجه بشه و چیزی رو که یوزر می‌خواهد رو تولید کنه. یعنی ساختار آموزش بجای «بقیه‌اش رو تو بگو» به ساختار ارباب رعیتی «این کاری که می‌گم رو بکن» تبدیل بشه. بطورکلی این دو نیاز ذیل تسک Alignment تعریف می‌شه و برای هم‌راستا کردن تولیدات مدل با ترجیحات انسانی تسک Alignment رو انجام می‌دن.

برخلاف تصور، ChatGPT مستقیماً از روی GPT-3 ایجاد نشده. بلکه از نظر OpenAI راه ChatGPT از fine-tune کردن InstructGPT می‌گذشته. در واقع GPT-3 رو با اصلاح ساختار آموزش و ارائه یک روش آموزشی نو ظهور، fine-tune کردند و InstructGPT به دنیا اومده. و بعد از این مدل به ChatGPT رسیدند. جالب اینجاست که اصل زیبایی‌های خلقت توی InstructGPT جمع شده. و از InstructGPT تا ChatGPT خیلی مسایل فنی خاصی رخ نداده.



برای ساخت InstructGPT در اولین مرحله اومدن روی مدل GPT-3 تسک Supervised fine tuning (SFT) رو انجام دادند. یعنی اومدن تمام Promptهایی که ملت روی GPT-3 داشتند رو به آدم‌ها دادند و ازشون خواستند پاسخ این Prompt رو بنویسند. و بعد با کمک این سوال و جواب‌ها مدل رو ذیل تسک SFT فاین تیون کردند. این اولین مرحله‌ی اون مکانیزمی است که در عنوان بهش اشاره کردم. یعنی مکانیزم Reinforcement Learning Human Feedback. بعد از انجام مرحله‌ی اول، مرحله‌ی دوم به این صورت انجام می‌شه:

مرحله دوم: اول به ازای هر Prompt، از مدلی که در مرحله‌ی اول فاین تیون شده چندین خروجی می‌گیریم. این خروجی‌ها را با تنظیمات مختلف می‌گیریم که خروجی‌های متفاوتی داشته باشیم. و یا در مرحله اول چندین مدل با معماری متفاوت فاین تیون می‌کنیم و در این مرحله از همه این مدل‌ها خروجی می‌گیریم تا خروجی‌های متفاوت داشته باشیم. خروجی‌ها رو به انسان می‌دیم تا برامون از بهترین تا بدترین جواب Sort کنه. (در اینجا Promptها، Requestهایی هستند که مردم روی مدل API GPT-3 زدند). و بعد از طریق این دیتای بارزش (ترتیب بندی نتایج مدل‌ها بر اساس ترجیح انسان) یک Reward Model توسعه می‌دیم. در واقع اینجا با این مدل داریم اون Functionی رو مدل می‌کنیم که معمولا یا Rule Based بود یا انسان.

مرحله سوم: در این مرحله می‌آیم مدل فاین تیون شده در مرحله‌ی اول رو تبدیل به یک مدل RL می‌کنیم. و به ازای هر Prompt در دیتابیس ازش خروجی می‌گیریم. خروجی رو میدیم به Reward Model و از Reward محاسبه شده برای آپدیت Policyهای مدل استفاده می‌کنیم.

حقیقتاً روش زیبایی است. بنظرم یک نکته مهم این روش در اینه که کاریدی و کار علمی-مهندسی در یک تعادل جذابی قرار داره. از یه طرف تبدیل کردن یه مدل زبانی به یک مدل RL بنظر خیلی خفن میاد و احتمالا بیشتر در آینده شاهدش باشیم. از طرفی، جایی که تصور نمی‌شد انسان حضور داشته باشه، از انسان استفاده شد (یعنی مرحله‌ی دوم). و در آخر هم با Reward Model زیبایی رو بر ما تمام کردند و در جایی که حضور انسان یا Rules پذیرفته شده بود اثبات کردند میشه مدلی ساخت که ترجیحات انسان‌ها رو مدل کرد و خلاصه که با RLHF نمایش زیبایی از تعامل انسان و ماشین رقم زدند.

در ChatGPT این مساله رو حل کردند که چطوری نحوه ارتباط کاربر با مدل حالت تعاملی بگیره (اصلاح پاسخ‌های قبلی و...) که بطور کلی تغییرات ریزی روی Reward Model دادند. ضمن اینکه چالش‌های Iterative Deployment رو برای تصحیح اشتباهات و همچنین ایمن کردن مدل از ورودی‌های سمی کاربران تاحدی حل کردند.

برای مطالعه بیشتر:

<https://openai.com/blog/instruction-following/#fn1>

<https://arxiv.org/abs/2203.02155>

<https://openai.com/blog/deep-reinforcement-learning-from-human-preferences>