



تراشه‌های هوش مصنوعی ۹۰ درصد ارزان‌تر می‌شوند

استارت‌آپ تراشه‌سازی تحت حمایت شرکت «مایکروسافت» می‌گوید ارزان شدن تراشه‌های هوش مصنوعی حتی در حال حاضر هم امکان‌پذیر است.

استارت‌آپ تراشه‌سازی تحت حمایت شرکت «مایکروسافت» می‌گوید ارزان شدن تراشه‌های هوش مصنوعی حتی در حال حاضر هم امکان‌پذیر است.

به گزارش آیسنا، جهان در حال حاضر به قدری شیفته هوش مصنوعی است که شرکت‌های نرم‌افزاری روی انرژی هسته‌ای سرمایه‌گذاری می‌کنند تا تقاضا را برای تولید متن، تصویر و ویدیو افزایش دهند. اما اگر نیازی به این کار نباشد، چه پیش خواهد آمد؟ چه می‌شود اگر بتوانیم هوش مصنوعی خودمان را با راندمان بالاتر داشته باشیم تا بتوانیم با هزینه کمتر و مصرف انرژی بسیار کمتر، کارهای بیشتری انجام دهیم؟

به نقل از فوربس، این ایده‌ای است که استارت‌آپ تراشه‌سازی «d-Matrix» با حمایت شرکت «مایکروسافت» (Microsoft) ارائه داده است. ایده اصلی، ساخت تراشه‌هایی است که استنتاج را بسیار سریع‌تر، ارزان‌تر و کارآمدتر ارائه می‌دهند. این کاری است که شرکت‌های هوش مصنوعی هنگام پاسخ دادن به پرسش‌های هوش مصنوعی شما انجام می‌دهند. اگر d-Matrix درست بگوید، آینده هوش مصنوعی ممکن است به این که چه کسی بزرگترین مدل‌ها را آموزش می‌دهد وابسته نباشد، بلکه احتمالاً به این بستگی خواهد داشت که چه کسی می‌تواند سریع‌ترین و ارزان‌ترین پاسخ را بدهد.

«سید شت» (Sid Sheth)، مدیرعامل d-Matrix اخیراً طی گفتگویی در «اجلاس وب قطر» به گزارشگر فوربس گفت: آموزش کاملاً عملکرد است و استنتاج کاملاً کارآیی.

این تمایز برای d-Matrix اساسی است. آموزش مدل‌های زبانی بزرگ امروزی، کار دشواری به شمار می‌رود که بهتر است روی پردازنده‌های گرافیکی رده بالای شرکت «انویديا» (Nvidia) یا TPU متعلق به «گوگل» یا تعداد انگشت‌شماری از تراشه‌های دیگر انجام شود اما شت معتقد است که پردازنده‌های گرافیکی برای اجرای مدل‌های هوش مصنوعی جهت پاسخ دادن به پرسش‌ها ایده‌آل نیستند. در هر حال، این دقیقاً همان چیزی است که صنعت از آن استفاده می‌کند؛ عمدتاً به این دلیل که این همان چیزی است که صنعت در اختیار دارد. این امر تقریباً به تمیز کردن خانه با چکش و میخ شباهت دارد فقط به این دلیل که از همین ابزار برای ساختن خانه استفاده شده است.

شت گفت: مشکل اساسی این است که شما از یک تراشه آموزشی استفاده می‌کنید و بعد می‌گویید: «من قرار است روی آن تراشه‌ها استنتاج را اجرا کنم.» این واقعاً بهترین راه نیست.

شرکت d-Matrix براساس این باور تأسیس شد که استنتاج در نهایت بر حجم کار هوش مصنوعی تسلط خواهد یافت. این شرکت به جای تغییر کاربری سخت‌افزار آموزشی، ساختار را از پایه بنا کرد. به گفته شت، یک تفاوت ساختاری اساسی بین تراشه‌های ساخته شده برای استنتاج و تراشه‌های ساخته شده برای آموزش وجود دارد. آموزش، یک مشکل محاسباتی است، اما به عقیده شت، استنتاج فقط یک مشکل محاسباتی نیست، بلکه یک مشکل محاسباتی و حافظه‌ای است. بخش حافظه، تأخیر را افزایش می‌دهد.

در مدل‌های زبانی بزرگ، مرحله اولیه پردازش اغلب «prefill» نامیده می‌شود. مدل زبانی بزرگ، اعلان را دریافت می‌کند، محتوا را می‌سازد و پارامترهای مرتبط را در حافظه بارگذاری می‌کند. پس از آن، تولید توکن‌ها را در مرحله «رمزگشایی» آغاز می‌کند تا پاسخ مورد نظر کاربر را ارائه دهد. این تمایز مهم است، زیرا عملکرد استنتاج نه تنها توسط محاسبات خام، بلکه توسط سرعت انتقال داده‌ها بین حافظه و واحدهای پردازش محدود می‌شود. اگر محاسبات سریع باشند، اما دسترسی به حافظه کند باشد، توکن‌ها متوقف می‌شوند. اگر حافظه سریع باشد، اما محاسبات نتوانند همگام شوند، توان عملیاتی کاهش می‌یابد. در هر صورت، کاربران با تأخیر روبه‌رو می‌شوند.

شت توضیح داد: مانند این است که من از شما یک سوال انتزاعی بپرسم. شما قرار نیست فقط یک پاسخ سریع بدهید. شما قرار است پردازش کنید. سپس، دو تا سه ثانیه بعد احتمالاً شروع به صحبت خواهید کرد. بخش صحبت کردن، رمزگشایی است. بخش عمده‌ای از زیرساخت‌های امروزی برای حجم کاری آموزشی تنظیم شده‌اند که عملکرد در اوج را بر پاسخگویی متوسط اولویت می‌دهند، اما در استنتاج به ویژه هنگام استفاده از هوش مصنوعی تعاملی با پرسش و پاسخ‌های متعدد، تأخیر به معیار تعیین‌کننده تبدیل می‌شود. شت اضافه کرد: کاربر انتظار دارد وقتی مدل شروع به پاسخ دادن کرد، با سرعت مشخصی پاسخ بدهد. در غیر این صورت، احتمالاً به آن گوش نخواهد داد.

گزارشگر فوربس از شت درباره سرعت موتورهای هوش مصنوعی برای ارائه پاسخ‌های خود در اپلیکیشن‌ها و مرورگرها پرسید که اغلب کند به نظر می‌رسد. پرسش این بود: «آیا این سرعتی است که آنها فکر می‌کنند ما می‌توانیم تحمل کنیم یا حداکثر سرعتی است که می‌توانند داشته باشند؟» پاسخ شت این بود: «حداکثر سرعتی است که می‌توانند داشته باشند.»

از نظر عملی، این به معنای به حداقل رساندن زمان لازم برای بازیابی و فعال‌سازی مدل‌ها و تغذیه آنها با واحدهای محاسباتی برای هر توکن تولیدشده است. ساختارهای سنتی GPU، محاسبات و حافظه با پهنای باند بالا را به عنوان زیرسیستم‌های مجزا از هم ارائه می‌دهند و این می‌تواند به ناکارآمدی در بار کاری استنتاج منجر شود که به شدت به حافظه نیاز دارد.

راه حل d-Matrix، ترکیب دقیق محاسبات و حافظه در معماری آن است. این شرکت با نزدیک‌تر کردن فیزیکی حافظه به محاسبات و تنظیم جریان داده‌ها به طور ویژه برای الگوهای استنتاج قصد دارد تأخیر توکن را کاهش دهد و تعداد توکن‌ها در ثانیه به ازای هر وات را بالا ببرد. به علاوه، این استارت‌آپ به جای ساخت یک پردازنده یکپارچه بزرگ، سیلیکون خود را به عناصر سازنده کوچک‌تر

و ماژولار به نام «چیپلت» (Chiplet) برش می دهد. چیپلت ها بسته به نیازهای حجم کار، در مقادیر گوناگون با یکدیگر ترکیب می شوند. از نظر مفهومی، این به طراحی حافظه یکپارچه اپل نزدیک تر است تا ساختارهای سنتی پردازنده های گرافیکی که کاهش فاصله بین محاسبات و حافظه و تنظیم حول محور بهره وری به جای عملکرد معیار در اوج را شامل می شود. استارت آپ d-Matrix، سیلیکون را با این فرض طراحی کرد که استنتاج غالب خواهد بود. شت ادعا می کند که نتیجه این کار، تأخیر کمتر و عملکرد بسیار بالاتر به ازای توان مصرفی به ویژه برای بارهای کاری تعاملی و بلادرنگ است. شت خاطرنشان کرد که هزینه اجرای عملیات استنتاج با d-Matrix در حال حاضر تقریباً ۹۰ درصد کمتر از انواع GPU است و موارد بیشتری هم در راه هستند. وی افزود: این متعلق به امروز است. از این بهتر هم خواهد شد. این وعده شاید درست به موقع یا شاید کمی دیرتر از راه برسد. شرکت «اوپن ای آی» (OpenAI) اکنون از تراشه های بزرگ شرکت الکترونیک آمریکایی «سریراس» (Cerebras) برای اجرای استنتاج در مدل «GPT-5.3-Codex-Spark» خود استفاده می کند و در مقایسه با سایر ساختارها به سرعت ۱۵ تا ۲۰ برابر دست می یابد. سریراس یک رویکرد کاملاً متفاوت را از d-Matrix دارد، اما اهداف مشابهی را دنبال می کند. شرکت d-Matrix در حال حاضر تراشه های خود را در مقادیر کم عرضه می کند. شت گفت که خیلی زود این تعداد به هزاران عدد خواهد رسید. این تعداد باید به زودی به میلیون ها عدد برسد و ممکن است امسال به این رقم دست یابد. وی افزود: امسال شاهد تولید انبوه آن خواهید بود.