



هدفون‌های هوش مصنوعی که صحبت چندین گوینده را هم‌زمان ترجمه می‌کنند

هدفون‌های مجهز به هوش مصنوعی می‌توانند ترجمه گروهی را به طور هم‌زمان با شبیه‌سازی صدا ارائه دهند.

هدفون‌های مجهز به هوش مصنوعی می‌توانند ترجمه گروهی را به طور هم‌زمان با شبیه‌سازی صدا ارائه دهند. به گزارش ایسنا، «توچائو چن» (Tuochao Chen) دانشجوی «دانشگاه واشنگتن» اخیراً از موزه ای در مکزیک بازدید کرد. چن اسپانیایی صحبت نمی‌کند. بنابراین، یک اپلیکیشن ترجمه را روی تلفن همراه خود اجرا کرد و میکروفون را به سمت راهنمای تور گرفت اما حتی در سکوت نسبی موزه، سر و صدای اطراف زیاد بود و متن حاصل فایده زیادی نداشت. به نقل از تک اکسپلور، اخیراً فناوری‌های گوناگونی ظهور کرده‌اند که ترجمه روان را نوید می‌دهند اما هیچ کدام از این فناوری‌ها مشکل چن را در فضاهای عمومی حل نکردند. برای مثال، عینک‌های جدید شرکت «متا» (Meta) فقط با یک بلندگوی مجزا کار می‌کنند. آنها پس از به پایان رسیدن صحبت گوینده، ترجمه صوتی خودکار را پخش می‌کنند. اکنون چن و گروهی از پژوهشگران دانشگاه واشنگتن یک سیستم هدفون طراحی کرده‌اند که هم زمان صحبت چندین گوینده را ترجمه می‌کند و در عین حال، جهت و کیفیت صدای افراد را حفظ می‌کند. این گروه پژوهشی، سیستم را با هدفون‌های نویزگیر موجود در بازار که به میکروفون مجهز هستند، ساخته‌اند. الگوریتم‌های این گروه پژوهشی، گویندگان متفاوت را در یک فضا جدا می‌کنند، آنها را هنگام حرکت دنبال می‌کنند، گفتار آنها را ترجمه می‌کنند و با تأخیر دو تا چهار ثانیه‌ای پخش می‌کنند. «شیام گولاکوتا» (Shyam Gollakota) استاد دانشکده علوم رایانه و مهندسی دانشگاه واشنگتن و پژوهشگر ارشد این پروژه گفت: کد دستگاه برای دیگران در دسترس است تا براساس آن کار کنند. سایر فناوری‌های ترجمه بر این فرض ساخته شده‌اند که فقط یک نفر صحبت می‌کند اما در دنیای واقعی نمی‌توانید فقط یک صدای رباتیک داشته باشید که برای چندین نفر در یک اتاق صحبت کند. ما برای اولین بار صدای هر شخص و جهتی را که صدا از آن می‌آید، حفظ کرده‌ایم. این سیستم سه نوآوری را در بر دارد. نخست این که وقتی روشن می‌شود، بلافاصله تشخیص می‌دهد چه تعداد اسپیکر در فضای داخلی یا خارجی وجود دارد. چن گفت: الگوریتم‌های ما کمی شبیه به رادار کار می‌کنند. بنابراین، آنها فضا را به صورت ۳۶۰ درجه مورد بررسی قرار می‌دهند و دائماً به روزرسانی می‌کنند تا مشخص شود چند نفر در حال صحبت کردن هستند. سپس سیستم، گفتار را ترجمه می‌کند و کیفیت بیان و بلندی صدای هر گوینده را هنگام اجرا روی یک دستگاه مجهز به تراشه «Apple M2» مانند لپ‌تاپ‌ها و هدست «اپل ویژن پرو» (Apple Vision Pro) حفظ می‌کند. این گروه پژوهشی به دلیل نگرانی‌های مربوط به حریم خصوصی پیرامون شبیه‌سازی صدا، از به کار بردن محاسبات ابری اجتناب کردند. در نهایت، هنگامی که گوینده‌ها سر خود را حرکت می‌دهند، سیستم همچنان به ردیابی جهت و کیفیت صدای آنها همراه با تغییرات صورت گرفته ادامه می‌دهد. این سیستم در ۱۰ محیط داخلی و خارجی آزمایش شد و در یک آزمایش با ۲۹ شرکت‌کننده، کاربران این سیستم را به مدل‌هایی که اسپیکرها را در فضا ردیابی نمی‌کردند، ترجیح دادند. در یک آزمایش جداگانه روی کاربران، بیشتر شرکت‌کنندگان تأخیر سه تا چهار ثانیه‌ای را ترجیح دادند زیرا سیستم هنگام ترجمه با تأخیر یک تا دو ثانیه‌ای، خطاهای بیشتری مرتکب می‌شد. این گروه پژوهشی در تلاش هستند تا سرعت ترجمه را در نسخه‌های آینده کاهش دهند. این سیستم در حال حاضر فقط روی گفتار روزمره کار می‌کند، نه زبان تخصصی مانند اصطلاحات فنی. پژوهشگران در این پروژه با زبان‌های اسپانیایی، آلمانی و فرانسوی کار کردند اما بررسی‌های پیشین روی مدل‌های ترجمه نشان داده‌اند که می‌توان آنها را برای ترجمه حدود ۱۰۰ زبان آموزش داد. چن گفت: این گامی به سوی از بین بردن موانع زبانی بین فرهنگ‌هاست. بنابراین، اگر من در خیابان مکزیک قدم بزنم، حتی اگر اسپانیایی صحبت نکنم هم می‌توانم صدای همه مردم را ترجمه کنم و بدانم چه کسی چه گفته است.