

تاکید دبیرکل مجمع جهانی تقریب مذاهب اسلامی بر اصول اخلاقی هوش مصنوعی



حجت الاسلام والمسلمین دکتر حمید شهریاری دبیرکل مجمع جهانی تقریب مذاهب اسلامی با سخنرانی در کنفرانس بین المللی علم، دین و هوش مصنوعی که در زاگرب کرواسی برگزار شد، بر اصول اخلاقی هوش مصنوعی تاکید کرد.

به گزارش روابط عمومی مجمع جهانی تقریب مذاهب اسلامی، حجت الاسلام والمسلمین دکتر حمید شهریاری دبیرکل مجمع گفت: بحث هوش مصنوعی امروز وارد زندگی روزمره ما شده است؛ جایی که تمامی فعالیت‌های ما اعم از کار، دانش، تفریحات و استراحت ما را شامل می‌شود. همه دانشگاهیان و تجار به این فکر افتاده‌اند که در شغل خود از آن بهره‌برداری کنند. عجیب آن است که هوش مصنوعی نیز اعلام آمادگی کرده که در همه ابعاد زندگی انسان وارد شود و بسته به گزینش هر شخص به او در انتخاب‌های کمک کند. امروز دیگر هوش مصنوعی یک فرضیه نیست بلکه یک همراه و همدم برای انسان در تمام لحظات زندگی است.

وی افزود: البته بحث هوش مصنوعی بسیار گسترده است و همه ابعاد حیات بشر را فراموش نمی‌کند و هر فردی باید براساس مسئولیت خودش به آن بپردازد. ما نیز از نگاه توصیه‌های اخلاقی به معنای عام آن وارد این بحث می‌شویم و نکاتی اخلاقی بیان می‌کنیم که باید مبنای تصمیم‌گیری آینده برای زیست بشر باشد. آینده‌ای که در کنار بهره‌مندی از فرصت‌های هوش مصنوعی، سواد هوشمندی‌سازی خود را ارتقا می‌دهد تا از آسیب‌های آن تا حد ممکن بر حذر باشد. ما باید چارچوبی برای کاربرد هوش مصنوعی طراحی کنیم تا به این دو هدف دست یابیم.

دبیرکل مجمع با اشاره به مزایا و معایب همه گیر شدن هوش مصنوعی گفت: هوش مصنوعی می‌تواند به عنوان ابزاری برای تقویت صلح و امنیت جهانی عمل کند، اما در عین حال، استفاده نادرست از آن می‌تواند تهدیداتی جدی برای ثبات بین‌المللی ایجاد کند. استفاده نظامی از هوش مصنوعی، مانند تسلیحات خودمختار، نگرانی‌هایی درباره تشدید درگیری‌ها و افزایش رقابت تسلیحاتی ایجاد کرده است. برای جلوگیری از این تهدیدات، لازم است چارچوب‌های قانونی بین‌المللی برای استفاده مسئولانه از هوش مصنوعی تدوین شود.

وی خاطرنشان کرد: از سوی دیگر هوش مصنوعی می‌تواند ابزارهایی برای پیش‌بینی درگیری‌ها، تحلیل داده‌های اقلیمی مرتبط با منازعات و مقابله با نفرت‌پراکنی فراهم آورد و با استفاده از هوش مصنوعی الگوهای پیش‌بینی درگیری‌ها را ترسیم کرد. آن‌ها نمونه‌هایی از برنامه‌هایی را ارائه می‌دهند که از مهارت‌های پیش‌بینی برای شناسایی مناطق پرخطر استفاده می‌کنند و به مسائل مربوط به داده‌های پایه و قابلیت اطمینان این ابزارها پرداخته‌اند. همچنین آنها چگونگی تقویت نفرت‌پراکنی در رسانه‌های اجتماعی توسط الگوریتم‌های هوش مصنوعی و استفاده از پردازش زبان طبیعی (NLP) برای مقابله با گسترش آن و چگونگی استفاده از یادگیری ماشین با داده‌های جغرافیایی و بصری برای مقابله با نقض حقوق بشر را تبیین کرده‌اند.

دکتر حمید شهریاری «مسئولیت‌پذیری، توجیه‌پذیری و اعتمادسازی» را از اصول حرفه‌ای به کارگیری هوش مصنوعی دانست و ادامه داد: این موارد نقش اساسی در تضمین استفاده ایمن، اخلاقی و قابل اعتماد از فناوری‌های هوش مصنوعی در تصمیم‌گیری‌های حساس دارد. این اصل بر نیاز به پاسخگویی انسانی تأکید می‌کند. مسئولیت‌پذیری در سیستم‌های هوش مصنوعی باید به صورت انسانی باقی بماند. با توجه به اینکه الگوریتم‌ها نمی‌توانند به طور معناداری پاسخگو باشند و مجازات شوند، لازم است انسان‌ها مسئولیت نتایج و پیامدهای تصمیمات اتخاذ شده را بر عهده بگیرند.

