

The European Union's Artificial Intelligence Act as the First Global Document in the Field of Governing the Online World

Mohammad Mehdi Seyed Naseri^{ID}

P.h.D. Medical Ethics and Law Research Center, Shahid Beheshti University of Medical Sciences, Tehran, Iran. sm.snaseri@gmail.com

Abstract

Artificial intelligence, one of the most prominent emerging technologies of nowadays, has significant potential to change and transform various dimensions of human life. While such novel technology offers new opportunities in various fields such as healthcare, transportation, education, etc., it also creates new major challenges for human rights. In its pioneering role in regulating artificial intelligence, the European Union drafted the first vivid proposal for an Artificial Intelligence Act in 2021. The aim of drafting it was to create a legal framework for the responsible development and use of artificial intelligence within the European Union. After about two years of negotiations, in December 2023, negotiators in the European Parliament and the Council of Europe finally reached a provisional agreement on the Artificial Intelligence Act, and at most, in February 2024, the Permanent Representatives Committee voted to approve the political agreement reached in 2023. Despite, many critics believe that the European Union's Artificial Intelligence Act has failed to protect human rights and has not considered fundamental human rights principles. The question stricken to the mind now is whether or not the European Union's Artificial Intelligence Act can be considered as a positive step towards regulating the well-use of artificial intelligence and the responsible use of this novel technology. The present research is applied in terms of purpose and descriptive-analytical in research type. Overall, it can be depicted that the European Union's Artificial Intelligence Act is a positive step towards regulating the use of artificial intelligence and the responsible hiring of the novel technology aforementioned. However, to ensure that law fully protects human rights, amendments must be made to it. Regarding measures such as limiting the scope of the national security exemption, increasing transparency obligations for law enforcement and immigration authorities, and making the list of high-risk systems more transparent, the European Union's Artificial Intelligence Act can become a truly powerful tool for protecting human rights in the age of artificial intelligence.

Keywords: Artificial Intelligence, European Union, Governance, Emerging Technologies, Children.

Cite this article: Seyed Naseri, M.M. (2024). The European Union's Artificial Intelligence Act as the First Global Document in the Field of Governing the Online World. *Philosophy of Law*, 3(2), p. 7-24.
<https://doi.org/10.22081/PHLQ.2025.70710.1083>

Received: 2024-04-19 ; **Revised:** 2024-06-05 ; **Accepted:** 2024-06-22 ; **Published online:** 2024-09-23

© The Author(s).

Article type: Research Article

Publisher: Baqir al-Olum University



<https://phlq.bou.ac.ir/>

قانون هوش مصنوعی اتحادیه اروپا به مثابه نخستین سند جهانی در زمینه حکمرانی بر جهان آنلاین

محمد مهدی سیدناصری 

دکتری، مرکز تحقیقات اخلاق و حقوق پزشکی، دانشگاه علوم پزشکی شهید بهشتی، تهران، ایران. sm.snaseri@gmail.com

چکیده

هوش مصنوعی یکی از برجسته‌ترین فناوری‌های نوظهور عصر حاضر، دارای پتانسیل قابل توجهی جهت تغییر و تحول ابعاد مختلف زندگی بشری است. این فناوری نوین، در حالی که فرصت‌های جدیدی را در زمینه‌های مختلف همچون مراقبت‌های بهداشتی، حمل و نقل، آموزش و... ارائه می‌دهد، ابر چالش‌های جدیدی را نیز در قبال حقوق بشر ایجاد می‌کند. اتحادیه اروپا، در راستای پیشگامی در تنظیم هوش مصنوعی، در سال ۲۰۲۱ نخستین پیش‌نویس قانون هوش مصنوعی را تدوین نمود. هدف از تدوین این قانون، ایجاد چارچوبی جهت توسعه و استفاده مسئولانه از هوش مصنوعی در حوزه اتحادیه اروپا بوده است. پس از حدود دو سال مذاکره، در دسامبر ۲۰۲۳ مذاکره کنندگان در پارلمان اروپا و شورای اروپا به توافق موقت در مورد قانون هوش مصنوعی دست یافتند و در نهایت فوریه ۲۰۲۴، کمیته نمایندگان دائم به تأیید توافق سیاسی حاصل شده در سال ۲۰۲۳ رأی داد. با وجود این، منتقدان زیادی معتقدند که قانون هوش مصنوعی اتحادیه اروپا در حمایت از حقوق بشر ناکام مانده و اصول اولیه حقوق بشر را در نظر نگرفته است. حال پرسش این است که آیا قانون هوش مصنوعی اتحادیه اروپا می‌تواند گامی مثبت در جهت قاعده‌مند ساختن استفاده از هوش مصنوعی و استفاده مسئولانه از این فناوری نوین محسوب گردد؟ پژوهش حاضر از نظر هدف، کاربردی و از نظر نوع پژوهش، توصیفی-تحلیلی است. در مجموع می‌توان گفت که قانون هوش مصنوعی اتحادیه اروپا گامی مثبت در جهت قاعده‌مند ساختن استفاده از هوش مصنوعی و استفاده مسئولانه از این فناوری نوین است. با این حال، برای اطمینان از اینکه این قانون به‌طور کامل از حقوق بشر محافظت می‌کند، می‌بایست اصلاحاتی در آن انجام شود. با اقداماتی همچون، محدود نمودن دامنه معافیت امنیت ملی، افزایش تعهدات شفاف‌ساز برای مقامات مجری قانون و نهادهای مهاجرتی و شفاف‌تر کردن فهرست سامانه‌های پرخطر، قانون هوش مصنوعی اتحادیه اروپا می‌تواند به ابزاری قدرتمند جهت محافظت از حقوق بشر در عصر هوش مصنوعی تبدیل گردد.

کلیدواژه‌ها: هوش مصنوعی، اتحادیه اروپا، حکمرانی، تکنولوژی‌های نوظهور، کودکان.

استاد به این مقاله: سید ناصری، محمد مهدی (۱۴۰۳). قانون هوش مصنوعی اتحادیه اروپا به مثابه نخستین سند جهانی در زمینه حکمرانی بر جهان آنلاین.

فلسفه حقوق، ۷۳(۲)، ص ۷-۲۴. <https://doi.org/10.22081/PHLQ.2025.70710.1083>

تاریخ دریافت: ۱۴۰۳/۰۱/۳۱؛ تاریخ اصلاح: ۱۴۰۳/۰۳/۱۶؛ تاریخ پذیرش: ۱۴۰۳/۰۴/۰۲؛ تاریخ انتشار آنلاین: ۱۴۰۳/۰۷/۰۲

ناشر: دانشگاه باقرالعلوم (ع)

نوع مقاله: پژوهشی

© نویسندگان



<https://phlq.bou.ac.ir/>

۱. مقدمه

در دنیای امروز، برخی به فکر افتاده‌اند با ربات‌ها همچون انسان برخورد کنند. در حال حاضر، ربات‌ها اجسام بی جانی تلقی می‌شوند که جزء اموال انسان‌ها هستند. در آینده، زمانی که ربات‌ها هوشمند شوند، شاید مجبور شویم راجع به این مسئله تجدیدنظر کنیم. آیا دادن حق و حقوق به ربات‌ها، واقعاً قریب الوقوع است؟ نقش هوش مصنوعی، هر روز در زندگی ما بیشتر و بیشتر می‌شود. شروع توسعه این تکنولوژی در واقع، به خیلی قبل‌تر برمی‌گردد؛ یعنی زمانی در دهه ۵۰ میلادی که دانشگاه دارتموث در ایالات متحده، یک پروژه تحقیقات تابستانی را به هوش مصنوعی اختصاص داد. حتی می‌توان ریشه‌های هوش مصنوعی^۱ را در عمق بیشتری از تاریخ و در فعالیت‌های «آلن نیوئل»، «هربرت ای. سیمون» و «آلن تورینگ» جست‌وجو کرد. آزمون مشهور تورینگ در دهه ۱۹۵۰ میلادی، توسط او در مقاله‌ای مطرح شد. این مقاله، یکی از اولین اسنادی است که در آن، به وجود آمدن ماشین‌های هوشمند پیش‌بینی شده است. با این حال، مقوله هوش مصنوعی تا پیش از معرفی شدن سوپر کامپیوتر «دیب بلو» توسط کمپانی آی بی ام^۲، کماکان توجه جهانیان را به خود جلب نکرده بود. الگوریتم‌های هوش مصنوعی برای سال‌های متمادی است که در دیتاسنترها و کامپیوترهای بزرگ استفاده می‌شوند؛ ولی حضور آنها در حوزه لوازم الکترونیک مصرفی به سال‌های اخیر برمی‌گردد. آلن تورینگ در دهه ۱۹۵۰ میلادی، جمله‌ای دارد که می‌گوید: پیشنهاد می‌کنم این پرسش بررسی شود که آیا ماشین‌ها می‌توانند فکر کنند؟ واقعیت امروزی هوش مصنوعی، آنچه تا امروز به آن دست یافته شده و آنچه ممکن است بعدها به آن دست یابند؛ بسیار هیجان‌انگیز است، اما با هوش مصنوعی در داستان‌های علمی و تخیلی، فاصله زیادی دارد. به نظر می‌رسد امنیت بین‌الملل و روابط بین‌الملل جاذبه‌های جدیدی برای ابتکارات هوش مصنوعی و برنامه‌های کاربردی است. در میان تمامی بخش‌های رفتاری انسان، احتمالاً سیاست، دشوارترین رفتاری است که بتوان آن را به صورت اتوماسیون درآورد. سیاست به‌طور ذاتی، امر بسیار پیچیده‌ای است که پیچیدگی رفتار انسان، هم به‌عنوان یک فرد و هم در ابعاد اجتماعی، را انعکاس می‌دهد. این پیچیدگی، در سطح روابط بین‌الملل بسیار واقعی‌تر به نظر می‌رسد.

در آینده نه‌چندان دور، با پدیدار شدن هوش مصنوعی در سطح انسانی، به نظر دور از ذهن است که بتوان آن را هوش مصنوعی عمومی نامید. حتی اگر پیشرفت در این زمینه، سریع‌تر از پیش‌بینی‌ها باشد؛ مقاومت قابل توجهی در برابر ایده تبدیل مسئولیت به ماشین وجود دارد. امروزه می‌توان این موضوع را به‌ویژه در مباحث مربوط به اتومبیل‌های بدون سرنشین و سیستم سلاح‌های رباتیک مشاهده کرد. شاید تا حدود ده سال پیش، استفاده از هوش مصنوعی و نقش آن در زندگی عادی ما، تنها محدود به مواردی

1. Artificial intelligence

2. IBM

محدود و بسیار فنی بود که جز اهل فن، توانایی فهم آن را نداشتند؛ اما حالا می‌توان گفت هوش مصنوعی جزئی از زندگی عادی ما انسان‌ها است. آیا ما می‌توانیم یک رئیس‌جمهور یا نخست‌وزیر رباتیک را در آینده تصور کنیم؟ چنین تصویری تا امروز دور از ذهن به نظر می‌آید، اما این بدان مفهوم نیست که هوش مصنوعی بر سیاست و روابط بین‌الملل، تأثیر قابل توجهی نخواهد گذاشت. این تأثیر، از طریق تغییر در روش‌های تصمیم‌گیری و یا اطلاعات انسان‌های تصمیم‌گیرنده، بروز و ظهور پیدا خواهد نمود؛ تا جایی که اجازه ندهند هوش مصنوعی تصمیم‌گیرنده باشد. ملاحظه کاربرد هوش مصنوعی در روابط بین‌الملل، شامل ساختارهایی است که تصمیم‌گیرندگان از آن حمایت می‌کنند. می‌توان اینگونه بیان کرد که سیستم‌های هوش مصنوعی، جایگزین انسان‌ها در سطوح بالای تصمیم‌گیری نخواهند شد؛ اما به‌طور فزاینده‌ای هوش مصنوعی بخشی از فضایی خواهد شد که تصمیم‌گیرندگان انسانی در آن فضا عمل می‌کنند. این سیر تکاملی، هم فرصت‌هایی بزرگ و هم خطرات قابل توجهی را به‌وجود می‌آورد. بنابراین، توجه به تأثیرات بالقوه در مراحل اولیه، بسیار حیاتی است. کاربرد هوش مصنوعی در خدمات حقوقی، به‌سرعت در حال افزایش است و سیستم‌های نوین مبتنی بر پردازش زبان طبیعی، به‌تدریج در حال به عهده گرفتن بخشی از وظایف حقوقدانان هستند. کارشناسان و متخصصان هوش مصنوعی، بر این اعتقاد هستند که در کشورهای قدرتمند و پیشرفته جهان، وکلا تا حدود ده سال آینده، شغل خود را از دست خواهند داد؛ زیرا، نرم‌افزارهای مبتنی بر تکنولوژی هوش مصنوعی در حال حاضر، امکان تهیه قراردادهای دقیق، تحلیل قراردادهای و اسناد حقوقی موجود و پیش‌بینی آرای دادگاه‌ها را فراهم کرده‌اند. هوش مصنوعی در نهایت، ممکن است قادر به اجرای تمامی وظایف عملیاتی یا شناختی که هوش انسانی در حال حاضر برای آن ضروری است، باشد. اما با توجه به این احتمال که چنین هوش مصنوعی دهه‌ها یا حتی قرن‌ها برای توسعه یافتن به زمان نیاز داشته باشد؛ ممکن است تحلیل‌گران و سیاستمداران کنونی، به‌طور منطقی روی وظایف اختصاص داده‌شده به هوش مصنوعی در کوتاه‌مدت تمرکز نمایند. اینچنین وظایفی، به‌شدت به قابلیت‌های هوش مصنوعی بستگی دارد. البته ماشین‌ها قادر به پردازش داده‌های بسیار زیاد و به‌صورت بسیار سریع می‌باشند. آنها همچنین می‌توانند حجم اطلاعات بسیار بالاتری را نسبت به ذهن انسان، در اختیار قرار بدهند. در حوزه قانونگذاری، می‌توان به‌طور خلاصه گفت، کنگره ایالات متحده آمریکا در سال ۲۰۲۲، کمیته‌ای را جهت بررسی ابعاد قانونی مسائل مرتبط با هوش مصنوعی تشکیل داده است. هوش مصنوعی، یکی از برجسته‌ترین فناوری‌های نوظهور عصر حاضر، دارای پتانسیل قابل توجهی برای تغییر و تحول ابعاد مختلف زندگی بشری است. این فناوری نوین، در حالی که فرصت‌های جدیدی را در زمینه‌های مختلف همچون مراقبت‌های بهداشتی، حمل و نقل، آموزش و... ارائه می‌دهد، چالش‌های جدیدی را نیز در قبال حقوق بشر ایجاد می‌کند. اتحادیه اروپا، در راستای پیشگامی در تنظیم هوش مصنوعی، در سال ۲۰۲۱ نخستین پیش‌نویس قانون هوش مصنوعی را تدوین نمود. هدف از

تدوین این قانون، ایجاد چارچوبی جهت توسعه و استفاده مسئولانه از هوش مصنوعی در حوزه اتحادیه اروپا بوده است. پس از حدود دو سال مذاکره، در تاریخ ۸ دسامبر ۲۰۲۳، مذاکره‌کنندگان در پارلمان اروپا و شورای اروپا به توافق موقت در مورد قانون هوش مصنوعی دست یافتند و در نهایت، در ۲ فوریه ۲۰۲۴، کمیته نمایندگان دائم، به تأیید توافق سیاسی حاصل شده در دسامبر ۲۰۲۳ رأی داد. مقررات قانون هوش مصنوعی، به منظور حفاظت از حقوق اساسی، دموکراسی، حاکمیت قانون و پایداری زیست محیطی در برابر هوش مصنوعی پرخطر تدوین شده است. در عین حال، قانون مزبور به دنبال تقویت نوآوری و تبدیل اروپا به رهبر و پیشگام در زمینه هوش مصنوعی است. با وجود این، منتقدان زیادی معتقدند که قانون هوش مصنوعی اتحادیه اروپا، در حمایت از حقوق بشر ناکام مانده و اصول اولیه حقوق بشر را در نظر نگرفته است. حال، پرسش این است که آیا قانون هوش مصنوعی اتحادیه اروپا، می‌تواند گامی مثبت در جهت قاعده‌مند ساختن استفاده از هوش مصنوعی و استفاده مسئولانه از این فناوری نوین محسوب گردد؟

۲. واکاوی مفهوم هوش مصنوعی

هوش، به شکل‌های مختلفی تعریف شده است. به‌طور کلی، می‌توان هوش را یک توانایی ذهنی بسیار کلی دانست که از جمله توانایی استدلال، برنامه‌ریزی، حل مسائل، تفکر انتزاعی، درک ایده‌های پیچیده، یادگیری سریع و یادگیری از طریق تجربه را شامل می‌شود. این توانایی، از صرف یادگیری یک کتاب یا یک مهارت تحصیلی محدود یا هوشمندی در آزمون، فراتر رفته و توانایی گسترده‌تر و عمیق‌تر را برای درک محیط پیرامون، منعکس می‌کند؛ همانند درک کردن و معنا کردن چیزها یا پیدا کردن آنچه باید انجام داد (Atluri, 2020, p. 15). هوش مصنوعی به قدرت استدلال و برنامه‌ریزی در رایانه یا سایر ماشین‌ها گفته می‌شود. هوش مصنوعی در مبنای، یکی از رشته‌های علوم رایانه است؛ که به مطالعه و توسعه دستگاه‌های هوشمند از طریق ارائه الگوریتم مناسب می‌پردازد؛ تا ماشین‌ها را قادر به ادراک، استدلال و یادگیری سازد. این رشته با تحقیق و توسعه نظریه‌ها، روش‌ها، فناوری‌ها و سیستم‌های کاربردی برای شبیه‌سازی و گسترش هوش انسانی، بر آن است که ماشین‌ها را قادر به انجام کارهای پیچیده‌ای کند که معمولاً برای انجام آنها به هوش انسانی نیاز است (بادینی و همکاران، ۱۳۹۳، ص ۱۱). البته ناگفته نماند که چنین ماشین‌هایی ممکن است هوش انسانی را تقلید، تقویت یا جایگزین کنند. هوش مصنوعی را می‌توان با قابلیت‌های خاصی دسته‌بندی کرد. به‌عنوان مثال، هوش مصنوعی ضعیف یا باریک، به هوش مصنوعی اطلاق می‌شود که می‌تواند رفتارهای هوشمند خاص انسان‌ها همچون تشخیص، یادگیری، استدلال و قضاوت را شبیه‌سازی کند. به عبارت دیگر، هدف هوش مصنوعی ضعیف، حل وظایف خاصی مانند تشخیص گفتار، تشخیص تصویر و ترجمه برخی از مطالب خاص است. هوش مصنوعی

قوی، به هوش مصنوعی اطلاق می‌شود که هوشیاری مستقل و توانایی نوآوری مشابه مغز انسان را دارد. هوش مصنوعی قوی می‌تواند فکر کند، برنامه‌ریزی کند و مشکلات را حل کند و همچنین درگیر تفکر انتزاعی شده، ایده‌های پیچیده را درک کرده، سریع یاد بگیرد و از تجربیات بیاموزد که در نتیجه، آن را به هوش انسانی نزدیک‌تر می‌کند. نوع دیگری از هوش مصنوعی، ابرهوش مصنوعی است که به هوش مصنوعی آینده اشاره دارد و از نظر توانایی محاسباتی و تفکر، بسیار از مغز انسان پیشی خواهد گرفت و بسیار باهوش‌تر از بهترین مغزهای انسان در هر زمینه‌ای؛ از جمله خلاقیت علمی، خرد عمومی و مهارت‌های اجتماعی خواهد بود (Batty, 2021, p. 21).

۳. چالش‌های کاربرد هوش مصنوعی در عصر انقلاب صنعتی چهارم

مانند سایر فناوری‌های نوظهور، هوش مصنوعی، یک شمشیر دو لبه است و با کاربرد نظامی گسترده‌تر هوش مصنوعی، مسائل جدیدی ظهور کرده که باعث گسترش نگرانی‌ها در سراسر جهان شده است. به عنوان مثال، با توجه به اینکه در حوزه مسائل نظامی، پتانسیل هوش مصنوعی در همه زمینه‌ها (به عنوان مثال زمین، دریا، هوا، فضا و اطلاعات) و همه سطوح جنگ (یعنی سیاسی، استراتژیک، عملیاتی و تاکتیکی) وجود دارد؛ یکی از چالش‌های استفاده از هوش مصنوعی در سطوح سیاسی و استراتژیک، آن است که این هوش ممکن است برای بی‌ثبات کردن حریف با تولید و انتشار مقادیر زیادی اطلاعات جعلی استفاده شود. از آنجا که بسیاری از تجهیزات مجهز به هوش مصنوعی، علاوه بر کارایی بالا به شفافیت بالا، ایمنی بالا و اعتماد یا درک کاربر نیاز دارند؛ به طور کلی، می‌توان مهم‌ترین چالش‌های موجود در کاربرد هوش مصنوعی در عرصه نظامی را مواردی همچون شفافیت، آسیب‌پذیری و یادگیری، دانست و همچنین الزاماتی که در سیستم‌های حیاتی، ایمنی، سیستم‌های نظارتی، عوامل مستقل، پزشکی و سایر کاربردهای مشابه نیز، معمول هستند. با پیشرفت‌های اخیر در توسعه هوش مصنوعی، علاقه تحقیقاتی به شفافیت، برای حمایت از کاربران نهایی در چنین برنامه‌هایی نیز افزایش یافته است. شفافیت، چالشی است که به دنبال حصول اطمینان از سازگاری عملکرد هوش مصنوعی با الزامات نظامی است. اگرچه از افزایش سرعت و دقت، می‌توان به عنوان مزایای بالقوه هوش مصنوعی در عرصه نظامی نام برد؛ اما همچنان نگرانی‌هایی در خصوص قابلیت‌هایی چون سرعت تصمیم‌گیری وجود دارد، که ممکن است سبب این امر شود که سیستم‌ها نتوانند با پیچیدگی‌های اجتناب‌ناپذیر جنگ، سازگار شوند. در نتیجه، امکان دارد چنین سیستم‌هایی نتوانند به طور دقیق بین رزمندگان و غیررزمندگان یا تهدیدها، تمایز قائل شوند و در نهایت، ممکن است دقت کمتری نسبت به اپراتورهای انسانی داشته باشند و اگر سیستم‌ها قبل از آزمایش کافی، فعال شوند یا اگر دشمنان موفق به جعل یا هک کردن آنها شوند؛ این مشکلات می‌توانند تشدید شوند (Emruli, 2016, p.18). این معیار، نشان می‌دهد که چگونه می‌توان

به طور بالقوه از جعل، استخراج داده‌ها و آلوده کردن داده‌های آموزشی جهت سوءاستفاده از آسیب‌پذیری‌ها و ایجاد تأثیرات منفی امنیتی نامطلوب استفاده کرد. آسیب‌پذیری‌ها و خطرات هوش مصنوعی، می‌تواند انواع مختلفی داشته باشد؛ از جمله آسیب‌پذیری‌های انسانی. به عنوان نمونه، بدون آموزش و بازآموزی نیروی کار جهت مطابقت با سرعت تغییرات فناوری و انواع مختلف تهدیدات، دشمنان هنگام تلاش جهت سوء استفاده از آسیب‌پذیری‌های هوش مصنوعی، با موانع کمتری مواجه خواهند شد. بنابراین، آموزش، خطاهای ناخواسته را کاهش می‌دهد.^۱ برای مثال، باتوجه به اینکه قابلیت‌های هوش مصنوعی دارای کاربردهای غیرنظامی و نظامی می‌باشد؛ سیستم‌هایی که به نیروهای نظامی در برنامه‌ریزی مأموریت و تخصیص منابع، کمک می‌کنند، ممکن است توسط مهاجمان برای شناسایی اهداف آسیب‌پذیر، مورد سوءاستفاده قرار گیرد (Gangjee, 2020, p. 28). توسعه برنامه‌های کاربردی مبتنی بر یادگیری ماشینی در زمینه نظامی، چالش‌برانگیز است؛ زیرا روش‌های جمع‌آوری داده‌ها در سازمان‌های نظامی، امکانات آموزشی، پلتفرم‌ها، شبکه‌های حسگر و سلاح‌ها در ابتدا، برای اهداف یادگیری ماشینی طراحی نشده بودند. در نتیجه، در این حوزه، یافتن مجموعه داده‌های واقعی با کیفیت و به اندازه کافی بزرگ، که بتوان از آنها برای یادگیری و کسب بینش استفاده کرد، اغلب دشوار است (Gongol, 2013, p.43).

۴. اصول اخلاقی حاکم بر کاربرد هوش مصنوعی

هوش مصنوعی، به سرعت در حال تبدیل شدن به بخشی از جنبه‌های زندگی در قرن بیست و یکم، از جمله جنگ است و با پیشرفته‌تر شدن و فراگیرتر شدن آن، یک پرسش اصلی، به طور فزاینده‌ای حیاتی جلوه می‌کند و آن، این است که آیا هوش مصنوعی می‌تواند اخلاقی باشد؟ بدیهی است که سیستم‌های مجهز به هوش مصنوعی نمی‌توانند ارزش‌های انسانی را درک کنند؛ امری که خود، موجب یکی از بزرگ‌ترین چالش‌های هوش مصنوعی از لحاظ اخلاقی است. در زمینه نظامی، این چالش اخلاقی بسیار عمیق‌تر شده و به ویژه مسائلی از قبیل نقض کرامت انسانی در مواجهه با سیستم‌های تسلیحاتی خودآیین را مطرح می‌کند (Grynberg, 2019, p. 7). باتوجه به اینکه عاملیت انسانی در هدایت مداخلات، صراحتاً در اسناد حقوق بشر دوستانه درج نشده است، تحقیق در مورد اخلاق و امنیت هوش مصنوعی، امری ضروری است. در این راستا، می‌بایست تلاش‌های مربوطه در حوزه فناوری و جامعه را یکپارچه نمود؛ تا اطمینان حاصل شود که توسعه هوش مصنوعی، همچنان برای انسان و طبیعت مفید است. بدیهی است فناوری و پیشرفت آن، الزامات جدیدی را برای کدهای اخلاقی ایجاد خواهد کرد؛ با این حال، با توجه به تفاوت‌های موجود از لحاظ فرهنگی و مکانی، هماهنگی استانداردهای اخلاقی میان دولت‌ها و

سازمان‌های مختلف بین‌المللی بسیار مهم است. در آوریل ۲۰۱۹، کمیسیون اروپا یک کد اخلاقی برای هوش مصنوعی، منتشر کرد و از راه‌اندازی مرحله آزمایشی این کد خبر داد و از شرکت‌ها و مؤسسات تحقیقاتی برای آزمایش آن دعوت کرد. در ۲۵ مه ۲۰۱۹، آکادمی هوش مصنوعی پکن، اصول هوش مصنوعی پکن را منتشر کرد.^۱ به موجب این اسناد، در مرحله تحقیق و توسعه، هوش مصنوعی بایستی تابع منافع کلی نوع بشر بوده و طراحی آن اخلاقی باشد. به منظور جلوگیری از سوءاستفاده در به‌کارگیری هوش مصنوعی، باید اطمینان حاصل شود که افراد ذینفع، از تأثیر بر حقوق و منافع خود، آگاهی و رضایت کامل داشته و از نظر حکمرانی، بایستی در مورد جایگزینی کار انسانی با هوش مصنوعی، با احتیاط عمل نمود. توصیه‌نامه یونسکو در خصوص هوش مصنوعی نیز «با در نظر داشتن این واقعیت که فناوری‌های مبتنی بر هوش مصنوعی می‌توانند خدمات مفید و قابل توجهی به زندگی بشر ارائه نمایند، به این حقیقت اذعان دارد که کاربرد بی‌قید و شرط این تکنولوژی، می‌تواند بنیان‌های مبانی اخلاقی و کرامت انسانی را به لرزه بيفکند» (Hoy, 2018, p. 3-4).

۵. قانون هوش مصنوعی اتحادیه اروپا به‌مثابه نخستین سند جهانی

در وهله نخست، این قانون، هوش مصنوعی را سامانه‌هایی مبتنی بر ماشین تعریف می‌کند که برای عملکرد با سطوح مختلفی از استقلال طراحی شده‌اند (به‌طور کامل مستقل یا تا میزانی با دخالت انسانی)؛ پس از استقرار، با گذشت زمان و با دریافت داده‌ها و اطلاعات جدید، توسعه و بهبود می‌یابند؛ قادرند از ورودی‌هایی که دریافت می‌کنند، استنباط و استنتاج داشته باشند^۲ و برای استخراج نتایج مختلف همچون پیش‌بینی‌های آینده، تولید محتوا، ارائه توصیه‌ها یا اتخاذ تصمیمات، استفاده شوند. قانون هوش مصنوعی، رویکردی مبتنی بر ریسک را برای محصولات و خدمات هوش مصنوعی اتخاذ می‌کند و تنها سامانه‌هایی را با سطوحی از خطرهای خاص، زیر چتر نظارتی خود قرار می‌دهد. سامانه‌های زیادی از جمله سامانه‌های توصیه محتوا و فیلترهای هرزنامه، اکثریت کم‌خطر را تشکیل می‌دهند. در مقابل، استفاده‌های پرخطر، مانند دستگاه‌های پزشکی و زیرساخت‌های حیاتی، با الزامات سخت‌گیرانه‌تری مواجه شده‌اند. علاوه بر این، برخی از کاربردهای هوش مصنوعی، همچون سیستم‌های امتیازدهی اجتماعی، پلیس پیشگیری و سیستم‌های تشخیص احساسات، به دلیل مخاطرات غیرموجه، به‌طور کامل، ممنوع شده‌اند (Roan & et al., 2023, p.23).

بدین ترتیب، قانون جدید هوش مصنوعی اتحادیه اروپا، الزامات متفاوتی را برای سیستم‌های کم‌خطر و پرخطر در نظر گرفته است. سیستم‌های کم‌خطر با الزامات کمتری نسبت به سیستم‌های پرخطر مواجه

1. <https://www.foley.com/en/insights/publications/2022/08/ai-regulation-where-china-eu-us-stand-today>

2. <https://www.smithsonianmag.com/smart-news/the-first-ai-lawyer-will-help-defendants-fight-speeding-tickets-180981508>

هستند. به عنوان مثال، سامانه‌های کم‌خطر نیازی به ارزیابی دقیق توسط نهادهای ناظر ندارند. ظاهراً، دلیل کم‌خطر شناخته شدن برخی سامانه‌ها، دسترسی نداشتن آنها به داده‌های شخصی حساس، مانند اطلاعات پزشکی یا مالی، نداشتن تأثیر چندان سوء و مستقیم بر زندگی افراد و قابلیت پایین آنها در سوءاستفاده برای نقض حقوق بشر است (Janssens, 2021, p.34). در مقابل، سامانه‌های پرخطر (مانند سیستم‌های تشخیص چهره، سیستم‌های امتیازدهی اجتماعی، سیستم‌های تشخیص احساسات، سیستم‌های هوش مصنوعی در حوزه‌های استخدام و مراقبت بهداشتی) که آن دسته از سیستم‌هایی هستند که می‌توانند خطرات قابل توجهی برای حقوق بشر و ایمنی افراد ایجاد کنند؛ مشمول الزاماتی همچون ارزیابی دقیق ریسک، مدیریت قوی داده، نظارت مستمر، شفافیت و پاسخگویی، واقع شده‌اند. بنابراین، تعهدات روشنی برای سامانه‌های هوش مصنوعی که به دلیل امکان ورود آسیب‌های چشمگیر به سلامت، ایمنی، حقوق اساسی، محیط زیست، دموکراسی و حاکمیت قانون، پرخطر تلقی می‌شوند، تعیین شده است. پیش از اینکه یک سیستم هوش مصنوعی پرخطر به کار گرفته شود، باید مورد ارزیابی قرار گیرد تا مشخص شود که آیا این سیستم به حقوق اساسی افراد، مانند حق حریم خصوصی، آزادی بیان یا برابری، آسیب می‌رساند یا خیر. همچنین، سامانه‌های هوش مصنوعی که برای تأثیرگذاری بر نتیجه انتخابات و رفتار رأی‌دهندگان طراحی می‌شوند، در زمره سیستم‌های پرخطر قرار گرفته‌اند. شهروندان، حق خواهند داشت که در مورد سامانه‌های هوش مصنوعی شکایت کنند و توضیحات لازم را درباره تصمیمات مبتنی بر سامانه‌های هوش مصنوعی پرخطر، که بر حقوق آنها تأثیر می‌گذارد، دریافت کنند.

همچنین، قانون‌گذاران اتحادیه اروپا درخصوص تهدیدات بالقوه هوش مصنوعی برای حقوق بشر، کاربردهای خاصی از هوش مصنوعی را ممنوع کرده‌اند. برای مثال، استفاده از سیستم‌های طبقه‌بندی بیومتریک که از ویژگی‌های حساس مانند عقاید سیاسی، مذهبی، فلسفی، گرایش جنسی و نژاد، برای شناسایی یا دسته‌بندی افراد استفاده می‌کنند، به دلیل نقض حریم خصوصی و تبعیض بالقوه علیه افراد، براساس ویژگی‌های ذاتی آنها ممنوع است. همچنین، جمع‌آوری و ذخیره‌سازی تصاویر صورت افراد از اینترنت یا فیلم‌های دوربین مدار بسته بدون هدف مشخص، مانند ایجاد پایگاه‌های اطلاعاتی تشخیص چهره، به دلیل نقض حریم خصوصی و سوءاستفاده احتمالی از اطلاعات جمع‌آوری شده، ممنوع شناخته شده است. استفاده از سیستم‌های تشخیص احساسات در محیط‌های کاری و مؤسسات آموزشی به دلیل نگرانی‌هایی که درخصوص نقض حریم خصوصی، سوءاستفاده از اطلاعات و تبعیض علیه افراد براساس حالات عاطفی آنها به دنبال دارد، کاربرد دیگری است که مشمول ممنوعیت این قانون، واقع شده است (Kur, 2014 p. 23).

گام مثبت دیگر در این قانون، ممنوعیت به‌کارگیری سامانه‌های رتبه‌بندی یا دسته‌بندی افراد بر پایه رفتار اجتماعی یا ویژگی‌های شخصی ایشان است؛ که به دلیل نقض حریم خصوصی، ایجاد تبعیض و

سوق دادن جامعه به سمت نظارت همگانی، مخاطره‌آمیز تلقی شده‌اند. ممنوعیت به‌کارگیری سامانه‌های هوش مصنوعی که با هدف دور زدن اراده‌ آزاد و دستکاری رفتار انسان طراحی شده‌اند؛ به دلیل نقض استقلال و آزادی انسان از موارد مهم دیگری است که باید بدان اشاره نمود. قانون‌گذاران همچنین، استفاده از هوش مصنوعی جهت سوءاستفاده از آسیب‌پذیری‌های افراد، اعم از سن، معلولیت، موقعیت اجتماعی یا اقتصادی را به دلیل نقض حقوق افراد و تشدید نابرابری‌های اجتماعی، ممنوع اعلام کرده‌اند. همه موارد پیش گفته، در راستای حفظ حقوق شهروندی و دموکراسی در عصر هوش مصنوعی، ضروری تلقی می‌شوند. به بیان دیگر، ممنوعیت‌های هوش مصنوعی، گامی مهم در جهت تضمین استفاده مسئولانه و اخلاقی از این فناوری نوظهور است. شایان ذکر است که، بر برخی از این ممنوعیت‌ها استثنائات و معافیت‌هایی نیز وارد شده است؛ ازجمله استفاده از سیستم‌های شناسایی بیومتریک^۱ به‌منظور جست‌وجوی هدفمند قربانیان (ربایش، قاچاق، بهره‌کشی جنسی)، جلوگیری از یک تهدید تروریستی خاص، یا شناسایی فردی مظنون به ارتکاب یکی از جرایم خاص ذکرشده در مقررات (مانند تروریسم، قاچاق، بهره‌کشی جنسی، قتل، آدم‌ربایی، تجاوز جنسی، سرقت مسلحانه، مشارکت در یک سازمان جنایی، جنایت زیست محیطی (Howel, 2022, p. 17).

علی‌رغم همه این موارد، نهادهای حقوق بشری ازجمله عفو بین‌الملل، قانون جدید هوش مصنوعی اتحادیه اروپا را به دلیل اولویت دادن به منافع صنعت و نهادهای مجری قانون به‌جای حقوق بشر، به‌شدت مورد انتقاد قرار داده‌اند. در حقیقت، این قانون به دلیل امتیازات ناشی از لابی‌گری در صنعت، معافیت‌های استفاده‌های خطرناک هوش مصنوعی توسط مقامات و نهادهای مجری قانون و مهاجرتی و خلأهایی که در برخی از خطرناک‌ترین فناوری‌های هوش مصنوعی ممنوع نشده است، مورد انتقاد قرار گرفته است. به بیان واضح‌تر، منتقدان معتقدند که این قانون به نفع شرکت‌های بزرگ فناوری است و به آنها اجازه می‌دهد تا در توسعه و استفاده از هوش مصنوعی نقش پررنگی داشته باشند. این امر، می‌تواند منجر به تسلط این شرکت‌ها بر بازار هوش مصنوعی و نادیده گرفته شدن حقوق و نیازهای مردم شود. اتحادیه اروپا، در این قانون، استفاده از هوش مصنوعی را برای اهداف امنیتی و مهاجرتی، از برخی الزامات قانونی معاف می‌کند. این امر، می‌تواند منجر به سوءاستفاده از هوش مصنوعی از سوی مقامات دولتی و نقض حقوق بشر شود. این قانون همچنین، برخی از خطرناک‌ترین فناوری‌های هوش مصنوعی، مانند سیستم‌های امتیازدهی اجتماعی و تشخیص احساسات را ممنوع نمی‌کند. منتقدان معتقدند که این کاستی می‌تواند منجر به استفاده از این فناوری‌ها برای اهداف نادرست و نقض حقوق بشر شود. در اوایل سال ۲۰۲۳، پارلمان اروپا طی اقدامی قابل‌توجه، ممنوعیت استفاده از تشخیص احساسات

را در چهار زمینه تصویب کرد: آموزش، محل کار، اجرای قانون و مهاجرت. این مصوبه، گامی مهم در جهت حفاظت از حقوق بشر و حریم خصوصی افراد در برابر سوءاستفاده از این فناوری نوظهور، به حساب می‌آید. اما متأسفانه، تحت فشار برخی از کشورهای عضو اتحادیه اروپا، ممنوعیت استفاده از تشخیص احساسات در زمینه‌های اجرای قانون و مهاجرت از متن نهایی قانون هوش مصنوعی اتحادیه اروپا، حذف شد. این امر، نشان‌دهنده رویکرد دوگانه و تبعیض‌آمیز این قانون در قبال حقوق اساسی افراد است. به‌طوری که افراد مهاجر و کسانی که در حاشیه جامعه قرار دارند، از حفاظت‌های قانونی کمتری برخوردار می‌شوند. حذف ممنوعیت استفاده از تشخیص احساسات در زمینه‌های اجرای قانون و مهاجرت، نگرانی‌های جدی را در مورد تبعات منفی این فناوری بر حقوق و آزادی‌های افراد ایجاد کرده است (Wills, 2022, p. 11). استفاده از این فناوری در این زمینه‌ها، می‌تواند به نقض حریم خصوصی، تبعیض و سوءاستفاده از افراد منجر شود؛ چرا که هنوز ابهامات زیادی در مورد نحوه قاعده‌مندسازی استفاده مقامات مجری قانون و نهادهای مهاجرتی از سامانه‌های هوش مصنوعی به کمک قانون هوش مصنوعی اتحادیه اروپا وجود دارد (Risse, 2021, p.14).

علاوه بر این، در حالی که ممنوعیت تشخیص احساسات در قانون هوش مصنوعی اتحادیه اروپا، گامی مثبت است؛ این قانون ظاهراً استثناهایی را برای اهداف پزشکی یا ایمنی، مجاز می‌داند. این خلأ می‌تواند بسیار خطرناک باشد؛ زیرا شرکت‌ها در حال حاضر، سامانه‌های به اصطلاح «تشخیص پرخاشگری» را می‌فروشند، که تصاویر مردان سیاه‌پوست را تهاجمی‌تر از مردان سفیدپوست می‌شناسند. استقرار چنین سیستم‌هایی در مدارس یا محل‌های کار، می‌تواند به نظارت نژادپرستانه از دانش‌آموزان یا کارگران سیاه‌پوست منجر شود. سیستم طبقه‌بندی پرخطر در قانون هوش مصنوعی اتحادیه اروپا، که مورد انتقاد Access Now (یک سازمان غیرانتفاعی که بر دفاع و گسترش حقوق دیجیتال کاربران در سراسر جهان تمرکز دارد) و دیگر سازمان‌ها بود، با اضافه شدن یک «فیلتر» به آن، بسیار ضعیف‌تر شد. پیشنهاد اولیه کمیسیون اروپا در این مورد، منطقی بود که همه موارد استفاده در فهرست برنامه‌های پرخطر، باید از تعهدات خاصی برای تضمین حقوق افراد پیروی می‌کردند (White & et al., 2020, p. 77). اما تحت فشار بخش صنعت و برخی دولت‌ها، این سیستم طبقه‌بندی با افزودن فیلتری که معیارهای آن مبهم و موسّع است، تغییر کرد. درواقع، با اضافه شدن این فیلتر، طراحان و توسعه‌دهندگان می‌توانند با اثبات این که سیستم هوش مصنوعی آنها در یکی از دسته‌بندی‌های خاص این فیلتر قرار می‌گیرد؛ از برخی یا کلیه تعهدات خود شانه خالی کنند (West & et al., 2018, p.4). در نتیجه، قانون هوش مصنوعی به‌طور واضح منافع صنعت را بر حقوق اساسی افراد در اولویت قرار می‌دهد. این نقص در قانون هوش مصنوعی، نگران‌کننده است و می‌تواند به نقض حقوق افراد در هنگام استفاده از سیستم‌های هوش مصنوعی پرخطر منجر گردد.

یکی از موارد بسیار مهمی که در مورد نگرانی‌ها نسبت به قانون هوش مصنوعی اتحادیه اروپا باید به آن اشاره کرد، معافیت امنیت ملی است که به دولت‌های عضو، اجازه می‌دهد تا در موارد خاص، از مقررات مربوط به هوش مصنوعی مستثنی شوند. هدف از این معافیت، حفظ توانایی کشورهای عضو در محافظت از امنیت ملی خود در برابر تهدیدات خارجی، شناخته شده است و شامل استفاده از سیستم‌های هوش مصنوعی برای اهداف نظامی، دفاعی و امنیت ملی، از جمله استفاده از سامانه‌های هوش مصنوعی برای جمع‌آوری اطلاعات و عملیات سرّی است. البته اعمال این معافیت، منوط به رعایت شرایطی است؛ از جمله، دولت‌های عضو باید دلایل موجهی برای استفاده از معافیت امنیت ملی ارائه دهند. استفاده از سیستم‌های هوش مصنوعی، باید ضروری و متناسب با تهدید امنیتی باشد و با حقوق و آزادی‌های اساسی افراد، تناسب داشته باشد و جبران‌های مؤثر و لازم در نظر گرفته شود. اما به هر حال، برخی از منتقدان این معافیت را بسیار گسترده می‌دانند و معتقدند که می‌تواند به سوءاستفاده از قدرت توسط دولت‌ها منجر شود و نگرانی‌هایی در مورد عدم شفافیت در مورد نحوه استفاده از این معافیت وجود دارد (Trappey & et al., 2019, p. 20-21).

۶. هوش مصنوعی و فلسفه اصل محافظت از کودکان

هوش مصنوعی قابلیت‌های گوناگونی را جهت افزایش امنیت سیستم و کاربر، غنی‌سازی مجموعه داده‌ها و پشتیبانی از مدل‌های تحلیلی بهبود یافته، فراهم می‌کند. هوش مصنوعی، حوزه وسیعی است که شامل یادگیری ماشینی و محاسبات شناختی می‌شود (Shnurenko & et al., 2020, p.12). در این حوزه، رایانه‌ها برای تقلید از عملکردهای شناختی انسان، مانند یادگیری و حل مسئله، برنامه‌ریزی می‌شوند؛ ولی به‌وضوح بسیار سریع‌تر و دقیق‌ترند. با الگوریتم‌های هوش مصنوعی، سیستم‌های محاسباتی، منطقی‌سازی می‌کنند و اقداماتی را با هدف دستیابی به غایتی خاص یا مجموعه‌ای از اهداف انجام می‌دهند. یکی از حُسن‌های آن، امنیت کاربران و ذی‌نفعان است؛ که می‌توان آن را از طریق ابزارهای مبتنی بر هوش مصنوعی، افزایش داد و راه‌های جدیدی را برای دسترسی به داده‌ها بدون مالکیت یا کنترل آن داده‌ها پدید آورد. درواقع، می‌توان از قابلیت‌های هوش مصنوعی در زمینه‌های مرتبط با حریم خصوصی بهره گرفت. هوش مصنوعی به منزله ابزاری مفید در محیط متاورس، با هدف حمایت مؤثر از کودکان استفاده می‌شود؛ برای نمونه، هوش مصنوعی، دسترسی به آموزش را برای کودکان با صرفه‌جویی در منابع (هم از نظر زمان و هم از نظر نیروی انسانی) ممکن می‌سازد. همچنین، به کمک هوش مصنوعی می‌توان فعالیت‌های متعارف در مقیاس بی‌سابقه انجام داد و راه‌های جدیدی را، نه فقط برای جمع‌آوری داده‌ها، بلکه برای پردازش آنها به‌منظور درک بهتر و ارزیابی نیازهای کودکان و ارائه خدمات مناسب‌تر به آنها پدید آورد. برای مثال، ترکیب کلان‌داده‌ها و هوش مصنوعی، امکان پردازش حجم وسیعی از داده‌های سلامت را

فراهم می‌کند (Seng, 2019, p. 8). این امر زمینه درمان هریک از بیماری‌های متعدد کودکان را فراهم می‌سازد و یکی از مهم‌ترین ابعاد سلامت زیستی کودکان است. در مقابل و با وجود اینکه استفاده از هوش مصنوعی در متاورس، برای کودکان فرصت‌های منحصر به فردی را فراهم می‌کند، یک نگرانی کلی وجود دارد که برنامه‌های مبتنی بر هوش مصنوعی، خطرات خاصی دارند که ممکن است به اندازه کافی از حقوق اساسی کودکان حمایت نکنند، یا حتی ممکن است در مواردی، این حقوق را نقض کنند. برای مثال، سیستم‌های مبتنی بر هوش مصنوعی، که اولویت‌های فردی کودکان را ثبت و محتوای شخصی‌سازی شده را ارائه می‌کنند، می‌توانند خطرات حریم خصوصی و داده‌ها را به همراه داشته باشند. چاره رفع این نگرانی، ارائه ابزارهای ویژه با هدف کاهش خطرات فناوری‌های نوظهور و استفاده از فرصت‌های بالقوه فناوری‌های مبتنی بر هوش مصنوعی برای کودکان است. (Sevastianova, 2020, p. 23)

در این خصوص، با وجود مواردی خاص، می‌بایست به طور ویژه در برابر تکنولوژی‌های نوین حساس بوده و بر چگونگی حمایت از کودکان در مقابل تکنولوژی‌های نوین، صراحت داشته باشند. مثالی مناسب برای توجه قانون‌گذار نسبت به ابزارهای حمایتی پیش‌گفته، نظام حقوقی اتحادیه اروپا است که در طول سال‌های متمادی از رویکرد توجه جزئی به این مسئله، به سمت سیاستی منسجم‌تر حرکت کرده است. در واقع از سال ۲۰۰۰، به موجب نظام حقوقی اتحادیه اروپا، حقوق کودکان که پیش‌تر در رابطه با زمینه‌های خاص بوده، به سمت مسیر جامع‌تری حرکت کرده است. برای نمونه، در خصوص خشونت و سوءاستفاده‌های آنلاینی که کودکان ممکن است تجربه کنند، اتحادیه اروپا در ماه می ۲۰۲۲، راهبردی جدید با نام اینترنت بهتر برای کودکان^۱ با هدف محافظت از کودکان و نوجوانان و تجهیز آنان به مهارت‌ها و ابزارهایی برای استفاده ایمن از اینترنت، اتخاذ کرده است. از منظر کمیسیون اتحادیه اروپا، این راهبرد جدید، مکمل راهبرد اتحادیه اروپا درباره حقوق کودک در سال ۲۰۲۱ میلادی بوده و منعکس‌کننده اصل محافظت از کودکان و جوانان در برابر فناوری‌های آنلاین است؛ تا بتوانند با این محیط، آگاهانه مواجه شوند. یکی از رویکردهای راهبرد جدید، این است که بتواند محیط دیجیتالی ایمن برای محافظت از کودکان در برابر محتوای آنلاین مضر و غیرقانونی در مقام مخاطب کم‌سن، فراهم کند. در واقع، فراهم کردن فضای آنلاین مناسب از طریق برنامه دیجیتالی امن و متناسب با سن که در راستای حفظ منافع کودکان باشد، هدف دیگر توانمندسازی دیجیتال است؛ تا همه کودکان، همچنین آنهایی که در موقعیت‌های آسیب‌پذیر قرار دارند، مهارت‌های لازم برای تعامل در محیط آنلاین را به صورت ایمن کسب کنند و یکی از مهم‌ترین سویه‌های سلامت زیستی در دنیای دیجیتال است (Ma & Baohong, 2020, p. 2-3).

در کنار آن، مشارکت فعال با انجام فعالیت‌های تحت کنترل کودکان برای کسب تجارب دیجیتالی

1. (BIK+) A European strategy for a better internet for kids (BIK+)

ایمن، نوآورانه و خلاقانه نیز، ازجمله رویکردهای این پیشنهاد است. همچنین ضرورت پیشگیری و مبارزه با سوءاستفاده جنسی از کودکان بهصورت آنلاین، ارائه‌دهندگان را موظف می‌کند که مطالب سوءاستفاده جنسی از کودکان را شناسایی، گزارش و مسدود و از خدمات خود حذف کنند و باید اقدامات لازم برای تأیید سن را نیز فراهم کنند (کیف پول هویت دیجیتال اروپایی، که کمیسیون پیشنهاد داده است، به تأیید سن کمک می‌کند). در نهایت، ارائه‌دهندگان خدمات دیجیتال باید شرایط و ضوابط خود را به‌گونه‌ای بنویسند که، کودکان بتوانند آن را درک کنند و از ارائه تبلیغات هدفمند براساس استفاده از داده‌های شخصی خردسالان جلوگیری کنند. همچنین، با بررسی موضوع در بسترهای حقوقی و قانونی ایران، روشن شد که به‌رغم اهمیت این امر و ضرورت توجه قانون‌گذار به مسئله توانمندسازی کودکان برای مواجهه با تکنولوژی‌های نوین، حمایت‌های خاص قانونی وجود ندارد.^۱ در این راستا، صرفاً اسنادی وجود دارند که به نظر می‌رسد مکفی و راهگشا نباشند. این اسناد، همچون مصوبه جلسه هفتادویکم شورای عالی فضای مجازی با موضوع «صیانت از کودکان و نوجوانان در فضای مجازی» است؛ که با وجود الزامات مناسب، توجه به تکنولوژی‌های نوین در آن کم‌رنگ است و موارد توجه به سلامت زیستی و کرامت انسانی در آن مورد توجه نمی‌باشد. جهت مواجهه هوشمندانه با زیست‌بوم متاورس، به‌طور کلی و حمایت مناسب از کودکان به‌طور خاص، باید متناسب با محیط متاورس، ابزارهایی خاص تعریف و استفاده شوند. نمونه‌ای قابل تأمل در این باره، چارچوب حریم خصوصی «ایکس. آر. اس. آی.» در حقوق آمریکا است (Rodrigues, 2020, p. 67). برنامه‌های واقعیت گسترده (ایکس. آر.)^۲ چندوجهی هستند و اغلب از مجموعه کامل حسگرهای موجود در یک دستگاه دارای ایکس. آر. خاص، استفاده می‌کنند. برای مثال، ابزارهای واقعیت گسترده، اغلب از داده‌های دوربین‌های تلفن همراه استفاده می‌کنند، که همراه با داده‌های سنسور شتاب‌سنج برای تعیین موقعیت دستگاه استفاده می‌شوند. هر مجموعه داده که جمع‌آوری می‌شود، با پیچیدگی‌هایی همراه است؛ از این‌رو، می‌بایست الزامات بیشتری لحاظ شود تا کنترل را هنگام پردازش و یا جمع‌آوری داده‌های مرتبط با خردسالان، تسهیل کند. ارائه‌دهندگان ایکس. آر. باید رویه‌ها و اقدامات امنیتی کافی را برای محافظت از داده‌های کودکان اجرا کنند؛ زیرا کودکان بخش گسترده‌ای از کاربران اصلی و آگاه به فناوری را تشکیل می‌دهند. همچنین، آنها هنگام جمع‌آوری و پردازش داده‌های شخصی از طریق کنترل‌کننده‌های داده، کمتر از خطرات موجود در پردازش داده‌های خود آگاه‌اند. ارائه‌دهندگان ایکس. آر. باید از همان ابتدای استفاده از فناوری‌های ایکس. آر.، از خردسالان محافظت کنند و این فناوری را با رعایت حریم خصوصی، طراحی و طبق اصول

1. <https://www.smithsonianmag.com/smart-news/the-first-ai-lawyer-will-help-defendants-fight-speeding-tickets-180981508>

2. Extended Reality (XR)

پیش فرض، پیاده سازی کنند. در مقام عمل، تشخیص اینکه آیا کاربر ایکس.آر، کودک است یا خیر؛ و برای مثال، رضایت معتبر والدین داده شده است یا خیر؛ اغلب دشوار است. بر این اساس، ارائه دهندگان ایکس آر، باید اقداماتی را که برای محافظت از داده های کودکان انجام می دهند، دائم بررسی کنند و می بایست بتوانند به جز تکیه بر مکانیسم های رضایت ساده، تأییدیه های مؤثرتری را اعمال کنند. به منظور تأمین الزامات پیش گفته، چارچوب حریم خصوصی «ایکس.آر.اس.آی»، ارائه شده است. این بستر براساس مقررات اروپایی حفاظت از داده «جی.دی.پی.آر.»، راهنمای مؤسسه ملی استانداردها و فناوری «ان.آی.اس.تی»، قانون حقوق آموزشی خانواده و حریم خصوصی «اف.ای.آر.پی.ای»، قانون حفاظت از حریم خصوصی آنلاین کودکان «سی.پی.پی.ای.»^۱ و چند قانون در حال تحول دیگر، تنظیم شده است. چارچوب حریم خصوصی «ایکس آر اس آی»، به سازمان ها کمک می کند تا اهداف حریم خصوصی خود را مشخص کنند، خطرات حریم خصوصی را شناسایی کرده و ضمن محدود کردن نقض حریم خصوصی، استفاده از اطلاعات شخصی و حساس را بهینه سازند. چارچوب حریم خصوصی «ایکس.آر.اس.آی»، از رویکرد چارچوب حریم خصوصی «ان.آی.اس.تی» الهام گرفته و از نظر راهبردی، به گونه ای طراحی شده که با نظام های حقوقی موجود در ایالات متحده و قواعد بین المللی، سازگار بوده و از طریق هر نوع سازمانی استفاده می شود؛ تا موجب پذیرش گسترده شود. هدف این چارچوب، اولویت قرار دادن حریم خصوصی و در نظر گرفتن آن در طول برنامه ریزی، طراحی، ساخت، استقرار، بهره برداری و از کار انداختن سیستم ها است. این چارچوب، منعطف است؛ برای مثال، سازمانی بزرگ که قبلاً برنامه حفظ حریم خصوصی قوی و فرآیندهای مدیریت ریسک صحیح داشته است، می تواند از این چارچوب برای تجزیه و تحلیل خطرات جدید حریم خصوصی و ایمنی استفاده کند. به همین ترتیب، سازمانی کوچک تا متوسط بدون برنامه حفظ حریم خصوصی نیز، می تواند از این چارچوب در حکم مرجعی برای درک و برآورده کردن انتظارات حریم خصوصی ذی نفعان خود استفاده کند. همچنین، با بررسی موضوع در نظام حقوقی ایران و تشخیص فقدان بسترهای مشابه، قانون گذار را ملزم به ارائه چارچوب های حریم خصوصی متناسب با تکنولوژی های نوین می کند. (Randakevi-Alpman, 2021, p.19)

۷. نتیجه گیری

هوش مصنوعی یکی از برجسته ترین فناوری های نوظهور عصر حاضر، دارای پتانسیل قابل توجهی جهت تغییر و تحول ابعاد مختلف زندگی بشری است. این فناوری نوین، در حالی که فرصت های جدیدی را در زمینه های مختلف همچون مراقبت های بهداشتی، حمل و نقل، آموزش و... ارائه می دهد، چالش های جدیدی را نیز در قبال حقوق بشر ایجاد می کند. اتحادیه اروپا، در راستای پیشگامی در تنظیم هوش

مصنوعی، در سال ۲۰۲۱ میلادی نخستین پیش‌نویس قانون هوش مصنوعی را تدوین کرد. هدف از تدوین این قانون، ایجاد چارچوبی جهت توسعه و استفاده مسئولانه از هوش مصنوعی در حوزه اتحادیه اروپا بوده است. پس از حدود دو سال مذاکره، در تاریخ هشتم دسامبر ۲۰۲۳ میلادی، مذاکره‌کنندگان در پارلمان اروپا و شورای اروپا به توافق موقت در مورد قانون هوش مصنوعی دست یافتند و در نهایت، در دوم فوریه ۲۰۲۴ میلادی، کمیته نمایندگان دائم به تأیید توافق سیاسی حاصل شده در دسامبر ۲۰۲۳ رأی داد. مقررات قانون هوش مصنوعی، به منظور حفاظت از حقوق اساسی، دموکراسی، حاکمیت قانون و پایداری زیست‌محیطی در برابر هوش مصنوعی پرخطر، تدوین شده است. در عین حال، قانون مزبور به دنبال تقویت نوآوری و تبدیل اروپا به رهبر و پیشگام، در زمینه هوش مصنوعی است. با وجود این، منتقدان زیادی معتقدند که قانون هوش مصنوعی اتحادیه اروپا در حمایت از حقوق بشر، ناکام مانده و اصول اولیه حقوق بشر را در نظر نگرفته است. قانون جدید هوش مصنوعی اتحادیه اروپا، شامل مجموعه‌ای از مقررات و محدودیت‌ها جهت استفاده از هوش مصنوعی است. برخی از این محدودیت‌ها عبارتند از:

شفافیت سیستم‌های هوش مصنوعی: این قانون، الزام شفافیت در سیستم‌های هوش مصنوعی را بیان می‌دارد. توسعه‌دهندگان و اپراتورهای هوش مصنوعی، می‌بایست اطلاعات واضح و قابل فهمی در مورد نحوه عملکرد سیستم‌های هوش مصنوعی، منطق پشت تصمیم‌هایشان و تأثیرات احتمالی این سیستم‌ها ارائه دهند.

مدیریت هوش مصنوعی پرخطر: این قانون، سیستم‌های هوش مصنوعی خاصی را به عنوان «پرخطر» شناسایی و طبقه‌بندی می‌کند که به نظارت دقیق‌تر نیاز دارند.

محدودیت در نظارت بیومتریک: این قانون، محدودیت‌های شدیدی را برای استفاده از فناوری‌های نظارت بیومتریک بلادرنگ، به ویژه در فضاهای در دسترس عموم، اعمال می‌کند.

محدودیت‌های کاربردی هوش مصنوعی: این قانون، برخی از برنامه‌های کاربردی هوش مصنوعی را که مضر یا با خطر بالا برای حقوق اساسی تلقی می‌شوند، ممنوع می‌کند.

چارچوب هوش مصنوعی با ریسک بالا: این قانون، چارچوب خاصی را برای سیستم‌های هوش مصنوعی ایجاد می‌کند که «خطر بالا» در نظر گرفته می‌شوند. قانون هوش مصنوعی اتحادیه اروپا، گامی مثبت در جهت قاعده‌مند ساختن استفاده از هوش مصنوعی و استفاده مسئولانه از این فناوری نوین است. با این حال، برای اطمینان از اینکه این قانون به طور کامل از حقوق بشر محافظت می‌کند، باید اصلاحاتی در آن انجام شود. با اقداماتی همچون، محدود کردن دامنه معافیت امنیت ملی، افزایش تعهدات شفاف‌ساز برای مقامات مجری قانون و نهادهای مهاجرتی، و شفاف‌تر کردن فهرست سامانه‌های پرخطر قانون هوش مصنوعی اتحادیه اروپا می‌تواند به ابزاری قدرتمند جهت محافظت از حقوق بشر در عصر هوش مصنوعی تبدیل شود.

منابع

- بادینی، حسن؛ حسین‌زاده، مجید؛ محبی‌فرد، سمانه (۱۳۹۳). بررسی نظریه استفاده منصفانه قانونی (کلاسیک) در علائم تجاری توصیفی. *پژوهشنامه بازرگانی*، ۱۹(۷۳)، ص. ۹۹-۱۲۴.
- Atluri, N. (2020). *A Study on Impact of Artificial Intelligence on Human Resource Management*. *Studies In Indian Place Names*, 40(58).
- Batty, R. (2021). Trade Mark Infringement & Artificial Intelligence. *New Zeal & Business Law Quarterly*, 26(3).
- Emruli, S. (2016). Trademark Protection, Absolute & Relative Grounds for Refusal of Trademark. *European Journal of Multidisciplinary Studies*, 192.
- Gangjee, D.S. (2020). *Eye, Robot: Artificial Intelligence & Trade Mark Registers*. In: N. Bruun, G. Dinwoodie, M. Levin & A. Ohly (eds.), *Transition & Coherence in Intellectual Property Law*. Cambridge University Press.
- Gongol, T. (2013). The preliminary ruling decision in the case of Google vs. Louis Vuitton concerning the adword service & its impact on the community law. *Amfiteatru Economic Journal*, 15(33).
- Grynberg, M. (2019). AI & the Death of Trademark. *Kentucky Law Journal*, 108(2).
- Howel, C. (2022). *AI Regulation: Where do China, the EU, and the U.S. Stand Today?* URL=<https://www.foley.com/en/insights/publications/2022/08/ai-regulation-where-china-eu-us-stand-today>
- Hoy, M. (2018). Alexa, Siri, Cortana & more: an introduction to voice assistants. *Medical reference services quarterly*, 37(1).
- Janssens, M.-Ch. (2021). *Challenges of AI Technology to Basic Notions of Trademark Law*. Una Europa Virtual Seminar: Mapping the Future of Law.
- Kur, A. (2014). Trademark's function, don't they? CJEU jurisprudence & unfair competition principles. *IIC-International Review of Intellectual Property & Competition Law*, 45(4).
- Ma, L. & Baohong, S. (2020). Machine learning & AI in marketing—Connecting computing power to human insights. *International Journal of Research in Marketing*, 37(3).
- Randakevi-Alpman, J. (2021). The Role of Trademarks on Online Retail Platforms: An EU Trademark Law Perspective. *GRUR International*, 70(7).
- Risse, M. (2021). The fourth generation of human rights: Epistemic rights in digital lifeworlds. *Moral Philosophy and Politics*, 8(2), p. 351-378.
- Roan, M., Young, L., Rudeliu, K. & Weissbrodt, D. (2023). *Study guide: freedom of religion or belief*. University of Minnesota Human Rights Center.
- Rodrigues, R. (2020). Legal And Human Rights Issues Of AI: Gaps, Challenges and Vulnerabilities. *Journal Of Responsible Technology*, 45(4).
- Sevastianova, V. (2020). *Trademark Functions in the Age of Voice Shopping: A Search Costs*

- Perspective*. Master's thesis. Department of Accounting & Commercial Law, Hanken School of Economics.
- Shnurenko, I., Murovana, T. & Kushchu, I. (2020). Artificial intelligence media and information literacy, human rights and freedom of expression. *Thenextminds for the unesco institute for information technologies in education*, no. 69.
- Trappey, A., Trappey, Ch. & Lin, S. (2019). Detecting trademark image infringement using convolutional neural networks. *Adv. Transdiscipl Eng.*, vol. 10.
- W. st, A., Clifford, J. & Atkinson, D. (2018). . lexa. . ui.. me.a brand'' An I. vestigat.on i. to th. impact of Artificial Intelligence on Branding. *The Business & Management Review*, 9(3).
- White, M., Mogyoros, A. & Gangjee, D.S. (2020). *IPO Artificial Intelligence & Intellectual Property: call for views-Trade Marks*. Oxford Intellectual Property Research Centre Faculty of Law University of Oxford.
- Wills, K. (2022). AI around the World: Intellectual Property Law Considerations & beyond. *Journal of the Patent & Trademark Office Society*, 102(2).

