

Ethical Requirements in the Approval Process of EU Artificial Intelligence Act

Reza Farajpour

PhD Student in Private Law, Faculty of Law, Islamic Azad University, Yasuj, Iran
(Corresponding Author) dr.r.farajpour@gmail.com

Mohamadbagher Amerinia

Associate Professor, Department of Law, Faculty of Law, Islamic Azad University,
Yasouj, Iran

Masoumeh Gorginia

PhD Student in Private Law, Faculty of Law, Islamic Azad University, Yasuj, Iran

DOI:10.71488/cyberlaw.2025.1127854

Keywords:

Ethics,
National Organization of
Artificial Intelligence,
Policymaking,
New Technology,
Law, Artificial Intelligence

Abstract

The development of artificial intelligence (AI) is rapidly changing the world. This technology has far-reaching effects on daily life, industries, and societies. With the rapid growth of artificial intelligence, ethical and legal issues related to the use of this technology are also increasing. The purpose of this research is to study the challenges and ethical and legal principles of artificial intelligence based on reviews of the domestic and foreign scientific sources in this field. The current qualitative research explains the legal and ethical challenges of artificial intelligence based on the analysis of the European Union's General Data Protection Regulation. The most important challenges in the functioning mechanism of artificial intelligence from an ethical point of view are privacy and data protection, transparency and accountability, justice and fairness, social and economic effects. Furthermore, from a legal point of view, violation of privacy and confidentiality, damage and responsibility for its compensation, institutional responsibility in cases of damage caused by the operation of the tools are regarded as the most important challenges. In this article, an attempt has been made to examine each of these cases and provide appropriate solutions. The use of ethical guidelines as a complementary law and governance tool in the development and use of artificial intelligence and analysis of this branch of science is considered very important, and that the developed countries are forming and enacting the related laws. The importance of artificial intelligence and the extent of its use in various sciences, including judicial sciences, have been so prominent that the European Union enacted a law entitled Artificial Intelligence Law on June 13, 2024. Therefore, it is necessary to take appropriate action in Iran with the help of prominent and renowned lawyers, experts in communication science and technology, and other related experts.



This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license: <http://creativecommons.org/licenses/by/4.0/>

الزامات اخلاقی در سیر تصویب قانون هوش مصنوعی اتحادیه اروپا

رضا فرج پور

دانشجوی دکتری حقوق خصوصی، دانشکده حقوق، دانشگاه آزاد اسلامی، یاسوج، ایران، (نویسنده مسئول)

پست الکترونیک: dr.r.farajpour@gmail.com

محمدباقر عامری نیا

دانشیار گروه حقوق، دانشکده حقوق، دانشگاه آزاد اسلامی، یاسوج، ایران

معصومه گرجی نیا

دانشجوی دکتری (حقوق خصوصی)، دانشکده حقوق، دانشگاه آزاد اسلامی، یاسوج، ایران

تاریخ پذیرش: ۱۸ شهریور ۱۴۰۳

تاریخ دریافت: ۱۱ مرداد ۱۴۰۳

چکیده

زمینه و هدف: توسعه هوش مصنوعی (AI) به سرعت در حال تغییر جهان است، و این فناوری تأثیرات گسترده‌ای بر زندگی روزمره، صنایع، و جوامع دارد. اما با رشد سریع هوش مصنوعی، مسائل اخلاقی و قانونی مرتبط با استفاده از این فناوری نیز به طور فزاینده‌ای مطرح می‌شود. هدف از پژوهش حاضر بررسی چالش‌ها و اصول اخلاقی و قانونی هوش مصنوعی می‌باشد و منابع علمی داخلی و خارجی را در این زمینه مرور می‌کند.

مواد و روش‌ها: پژوهش کیفی حاضر به تبیین چالش‌های حقوقی و اخلاقی عملکرد هوش مصنوعی بر اساس تحلیل آیین‌نامه عمومی حفاظت از داده‌های خصوصی اتحادیه اروپا می‌پردازد.

یافته‌ها: مهم‌ترین چالش‌های موجود در سازوکار عملکرد هوش مصنوعی از نظر اخلاقی عبارت است از حریم خصوصی و حفاظت از داده‌ها، شفافیت و پاسخگویی، عدالت و انصاف، تأثیرات اجتماعی و اقتصادی.

نتیجه‌گیری: استفاده از دستورالعمل‌های اخلاقی به عنوان تکمیل‌کننده قانون و ابزار حاکمیتی در توسعه و استفاده از هوش مصنوعی و تجزیه و تحلیل این شاخه از علم بسیار مهم تلقی می‌گردد و اغلب کشورهای توسعه یافته در حال نگارش و وضع قوانین مرتبط هستند. اهمیت هوش مصنوعی و گستردگی استفاده از آن در علوم مختلف از جمله علوم قضایی به حدی پررنگ بوده که اتحادیه اروپا قانونی با عنوان قانون هوش مصنوعی را در تاریخ ۱۳ ژوئن ۲۰۲۴ وضع نمود. لذا لازم است که در کشور ایران هم با استفاده از حقوق‌دانان برجسته و صاحب سبک و متخصصین علوم و فناوری ارتباطات و سایر صاحب نظران مرتبط نسبت به این مهم اقدام شایسته انجام پذیرد.

واژگان کلیدی: اخلاق، سازمان ملی هوش مصنوعی، سیاست‌گذاری، فناوری نوین، قانون، هوش مصنوعی

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 مانند 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 کرد. هدف از تدوین این قانون، ایجاد چارچوبی برای 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.

هوش مصنوعی به عنوان یکی از مهم‌ترین پیشرفت‌های فناوری، قابلیت‌های بی‌نظیری را در زمینه‌های مختلف ایجاد کرده است. از تشخیص بیماری‌ها تا بهینه‌سازی فرایندهای صنعتی، هوش مصنوعی توانسته است تغییرات چشمگیری را به ارمغان بیاورد. با این حال، این فناوری می‌تواند موجب نگرانی‌هایی در زمینه‌های اخلاقی و قانونی شود که نیازمند بررسی و مدیریت دقیق است. هوش مصنوعی با همه پتانسیل‌ها و فرصت‌هایی که به ارمغان می‌آورد، چالش‌های اخلاقی خاص خود را نیز دارد. رعایت الزامات اخلاقی در طراحی و استفاده از این فناوری نه تنها به جلوگیری از پیامدهای منفی کمک می‌کند، بلکه باعث افزایش اعتماد و پذیرش عمومی نسبت به هوش مصنوعی می‌شود. با توجه به سرعت پیشرفت تکنولوژی، لازم است که این الزامات اخلاقی همواره مورد بازنگری و به‌روزرسانی قرار گیرند تا همگام با تحولات جدید، بتوان بهترین استفاده را از هوش مصنوعی برد. امروزه هوش مصنوعی به سرعت در حال تحول و گسترش است و در بسیاری از جنبه‌های زندگی انسان‌ها تأثیرگذار شده است، از خودروهای خودران گرفته تا تشخیص‌های پزشکی، هوش مصنوعی نقش کلیدی در پیشرفت تکنولوژی دارد. با این حال، استفاده از این فناوری پیشرفته نیازمند رعایت اصول و الزامات اخلاقی خاصی است تا از پیامدهای منفی و سوءاستفاده‌های احتمالی جلوگیری شود. در این مقاله به بررسی مهم‌ترین الزامات اخلاقی در کاربرد هوش مصنوعی می‌پردازیم (رهبری و شعبانپور، ۱۴۰۱: ۴۱۹).

در فرآیند تصویب قانون هوش مصنوعی اتحادیه اروپا، چندین مسأله اخلاقی مطرح می‌شود، از جمله:

- حریم خصوصی و حفاظت از داده‌ها: با توجه به استفاده گسترده از داده‌های شخصی در سیستم‌های هوش مصنوعی، چگونه می‌توان از حریم خصوصی افراد محافظت کرد و تضمین نمود که داده‌ها به‌صورت اخلاقی و قانونی استفاده می‌شوند؟
- شفافیت و پاسخگویی: سیستم‌های هوش مصنوعی باید قابل فهم و شفاف باشند. چگونه می‌توان تضمین کرد که این سیستم‌ها به‌طور عادلانه و بدون تبعیض عمل کنند و مسئولیت‌پذیر باشند؟

- عدالت و انصاف: در طراحی و پیاده‌سازی هوش مصنوعی، چگونه می‌توان از تبعیض و نابرابری جلوگیری کرد و اطمینان حاصل کرد که این سیستم‌ها برای همه افراد به‌طور عادلانه عمل می‌کنند؟
- تأثیرات اجتماعی و اقتصادی: توسعه و کاربرد هوش مصنوعی می‌تواند تأثیرات عمیقی بر بازار کار، اقتصاد، و ساختارهای اجتماعی داشته باشد. چگونه می‌توان این تأثیرات را به‌طور اخلاقی مدیریت کرد؟

۱- پیشنهاد تحقیق

علی‌رغم توجه فزاینده‌ای که هوش مصنوعی و روش‌های توسعه یافته آن در تحقیقات چند رشته‌ای، رسانه‌ها و سیاست‌گذاری‌ها به خود جلب می‌کند، توافق روشنی در مورد اینکه چگونه هوش مصنوعی باید به بهترین شکل تعریف شود، وجود ندارد. به نظر می‌رسد که این موضوع نه تنها با توجه به برداشت عمومی، بلکه به علم کامپیوتر و قانون مرتبط است به عنوان مثال، گاسر و آلمیدا (Gasser & Almeida, 2017: 59) معتقدند که یکی از دلایل دشواری تعریف هوش مصنوعی از منظر فنی این است که هوش مصنوعی یک فناوری واحد نیست، بلکه مجموعه‌ای از تکنیک‌ها و زیرشاخه‌ها از حوزه‌هایی مانند تشخیص گفتار و بینایی رایانه گرفته تا حافظه و دقت توجه است. گروه تخصصی هوش مصنوعی کمیسیون اتحادیه اروپا هوش مصنوعی را چنین تعریف نمودند:

هوش مصنوعی به سیستم‌هایی اطلاق می‌شود که با تجزیه و تحلیل محیط خود و انجام اقدامات با درجاتی از خودمختاری برای دستیابی به اهداف خاص، رفتار هوشمندانه‌ای را نشان می‌دهند سیستم‌های مبتنی بر هوش مصنوعی می‌توانند صرفاً مبتنی بر نرم‌افزار باشند و در دنیای مجازی عمل کنند (مانند دستیارهای صوتی، نرم‌افزار تحلیل تصویر، موتورهای جستجو، سیستم‌های تشخیص گفتار و چهره) یا هوش مصنوعی را می‌توان در دستگاه‌های سخت افزاری تعبیه کرد (مانند روبات‌های پیشرفته، اتومبیل‌های خودران، پهپادها یا برنامه‌های کاربردی اینترنت اشیا). سؤالی که ذهن من را به عنوان یک پژوهشگر به خود معطوف نموده این است که آیا همیشه هوش مصنوعی رفتار هوشمندانه و منطقی دارد؟

پیچیدگی ساختار مفهومی باعث شد که گروه کارشناسان خبره مجدداً تعریفی نسبتاً پیچیده ارائه دهند که اولین تعریف کمیسیون اتحادیه اروپا را گسترش می‌دهد که بیان می‌نماید: سیستم‌های هوش مصنوعی سیستم‌های نرم‌افزاری و سخت‌افزاری طراحی شده توسط انسان که با توجه به یک هدف پیچیده، در بعد فیزیکی یا دیجیتالی با درک محیط خود از طریق جمع‌آوری داده‌ها، تفسیر داده‌های ساختار یافته یا بدون ساختار جمع‌آوری شده و استدلال در مورد آنها عمل می‌کنند. بنابراین، جنبه‌های متفاوتی از هوش مصنوعی وجود دارد که باید در تعریف هوش مصنوعی به عنوان چالشی در برابر مقررات در نظر گرفته شود. (Lipton, 2018: 33)

به نظر من این تعریف دارای پارادوکس است چراکه بر استقلال متمرکز است یعنی معیاری از عاملیت در سیستم‌های هوش مصنوعی وجود دارد و اشاره می‌کند که این سیستم‌ها می‌توانند هم از روبات‌های فیزیکی و هم از سیستم‌های نرم‌افزاری خالص تشکیل شوند در عین حال آنچه که یک ربات «پیشرفته» را مشخص می‌کند لزوماً مستلزم یک مرزبندی ساده نیست، که می‌توانیم انتظار داشته باشیم در طول زمان تغییر کند. به نظر اینجانب که حدود ۶ سال است در زمینه حقوق هوش مصنوعی و حقوق ربات‌ها در حال تحقیق و پژوهش بوده‌ام و با اساتید مطرح دانشگاه‌های آمریکا من جمله پروفیسور دیوید جی گانکل در ارتباط بوده و حتی کتاب‌های او را ترجمه نموده‌ام هوش مصنوعی را چنین تعریف می‌نمایم: برنامه‌ریزی و ارائه الگوریتم مناسب از سوی انسان به رایانه و درخواست از او جهت سهولت، تسریع و دقت در فرآیند جایگزینی هوش

انسانی به هوش مصنوعی. در تمامی مراحل ارائه خدمات هوش مصنوعی به انسان، باید درک منطقی، استدلال منطقی و یادگیری منطقی وجود داشته باشد (جی گانکل، ۱۴۰۱: ۱۲).

۲- چالش‌های حقوقی

چالش‌های حقوقی پیرامون عملکرد هوش مصنوعی در حوزه سلامت در موضوعات نقض حریم خصوصی و محرمانگی، خسارت و مسئولیت جبران آن و مسئولیت‌پذیری نهادی به شرح ذیل خلاصه می‌شوند:

۲-۱- نقض حریم خصوصی و محرمانگی:

پردازش داده‌های اشخاص، ارتباط تنگاتنگ با حفظ حریم خصوصی اشخاص دارد. به همین جهت است که مواد ۶ و ۱۲ آیین‌نامه بر ضرورت آگاهی بخشی به دارنده اطلاعات در خصوص کمیت و کیفیت اطلاعات مورد پردازش وی اشاره دارد، اما مسأله آگاهی دارنده اطلاعات از داده‌های پردازش شده یکی از چالش‌هایی است که بررسی ابعاد مختلف آن حاکی از نقص مقررات مصوب ۲۰۱۶ دارد. بند اول از ماده ۲۰ این مقررات بر ضرورت ارائه اطلاعات پردازش شده به دارنده اشاره دارد، اما سؤال پیش‌رو این است که دارنده اطلاعات آیا امکان بازخوانی و تحلیل داده‌های پردازش شده را خواهد داشت؟ آیا نهاد حکومتی در زمینه ارائه آموزش‌های لازم در این زمینه وجود داشته و چه آموزش‌هایی به آحاد جامعه در این زمینه باید ارائه گردد؟ نقش حکومت در فرایند نظارت بر دریافت و ارائه اطلاعات پردازش شده از سوی کنترل‌کننده به دارنده چه بوده و آیا این اطلاعات باید توسط حکومت نیز مورد بازبینی واقع شوند؟

مسأله دیگر در زمینه آگاهی بخشی به دارنده اطلاعات که ارتباط تنگاتنگ با حریم خصوصی وی دارد، وجود رضایت وی در پردازش داده‌های شخصی می‌باشد. بند اول از ماده ۶ آیین‌نامه مصوب ۲۰۱۶ بر ضرورت کسب رضایت موردی دارنده در پردازش داده‌های خصوصی وی اشاره دارد. مطابق با مفاد این مقرره نهادهای کنترل‌کننده در جمع‌آوری و انتقال اطلاعات شخصی اشخاص، باید رضایت آنها را در خصوص کمیت و کیفیت داده‌های مورد پردازش کسب نمایند. اما سؤال مهم این است که آیا آحاد یک جامعه آگاهی چندانی از فرایند پردازش اطلاعات خود و ضرورت یا عدم ضرورت پردازش داده‌های مذکور برخوردار می‌باشند؟ آیا محول نمودن وظیفه آگاهی بخشی به یک کنترل‌کننده هوش مصنوعی که خود ذی‌نفع عدم آگاهی هرچه بیشتر دارنده در فرایند پردازش می‌باشد، می‌تواند اهداف ناشی از تصویب مقررات ۲۰۱۶ را تأمین نماید؟ چالش مذکور زمانی جلوه بیشتر می‌یابد که کنترل‌کننده هوش مصنوعی، جهت پردازش داده نیاز به ارسال اطلاعات جمع‌آوری شده به یک شرکت فراملی پردازنده یا دارای تابعیت کشوری ثالث داشته باشد. در این خصوص مسأله ضرورت پردازش داده جهت انجام وظایف از پیش تعیین شده ابزار از یک طرف و از طرف دیگر ضرورت حفظ امنیت داده‌های مورد پردازش مسائلی است که محول نمودن تعیین تکلیف این امر بر تصمیم یک دارنده صرف، می‌تواند در مواردی خطرات جبران‌ناپذیری را فراهم آورد چراکه هرگونه سوءاستفاده از داده‌های شخصی اشخاص از جمله اطلاعات زیستی آنها می‌تواند زمینه ایجاد سلاح‌های زیستی و یا تولید ویروس‌های آزمایشگاهی و به خطراتدن امنیت ملی یک کشور را فراهم آورد.

مشکل دیگر مفاد بند سوم از ماده ۲۰ این مقررات می‌باشد. این بند تبادل و پردازش اطلاعات اشخاص در راستای حفظ منافع عمومی را بدون کسب رضایت دارنده میسر نموده است. عدم پیش‌بینی ابعاد دقیق عبارت «منافع عمومی» می‌تواند زمینه نقض شدید حریم خصوصی افراد را دربرداشته باشد، چراکه یک کنترل‌کننده یا پردازنده امکان تفسیر هر یک از اعمال خود در جهت حفظ منافع عمومی را برخوردار بوده و معیار و شرایطی برای این موضوع در مقررات مرقوم پیش‌بینی نشده

است. سوءاستفاده از مقرر شده بیان شده نه تنها می تواند زمینه نقض وظایف مرتبط با آگاهی بخشی به دارنده و کسب رضایت او را فراهم آورد، بلکه زمینه سوءاستفاده و پردازش پنهانی اطلاعات وی را نیز می تواند دربرداشته باشد، به خصوص این که در این مقررات وظیفه ای نیز برای دولت متبوع دارنده در جهت کیفیت نظارت بر اجرای مقررات این بند پیش بینی نشده است.

۲-۲- خسارت و مسئولیت جبران آن

راه حلی که در راستای حل معضلات بیان شده می توان بیان داشت، ضبط اطلاعات و بررسی آنها به وسیله انسان می باشد تا در صورت وقوع خطا در عملکرد هوش مصنوعی قابلیت پیشگیری از آن داشته باشد. اما این امر به منزله در اختیار گرفتن وجود اطلاعات شخصی اشخاص توسط انسان می باشد که با مقررات حفاظت از اطلاعات افراد در اتحادیه اروپا واجد تعارض است. مطابق با ماده ۶ آیین نامه مصوب ۲۰۱۶ دسترسی به اطلاعات شخصی اروپاییان باید با رضایت موردی آنها باشد. در صورتی که چنین عملکردی از سوی نهاد عامل صورت نگرفته باشد در صورت وقوع هرگونه خسارت نهاد عامل مسئول جبران خسارات وارده است. مطابق با بند اول از ماده ۴ آیین نامه، اطلاعات شخصی افراد به هرگونه اطلاعاتی بیان می گردد که عرفاً از قابلیت اشتراک به عموم جامعه برخوردار نباشد. مطابق با بند یازدهم از ماده مرقوم و مواد ۶ و ۱۲ آیین نامه، کسب رضایت افراد در دسترسی و پردازش اطلاعات وی باید به صورت خاص و بدون ابهام توسط عامل صورت پذیرد. از این رو در صورتی که فرد به هر شکل در خصوص اطلاعات به دست آمده توسط وی آگاهی و رضایت نداشته باشد، مسئولیت های مندرج در آیین نامه مرقوم شامل حال وی می گردد.

این موضوع در نظام حقوقی ایران نیز موکداً در مواد ۵۹ و ۵۸ قانون تجارت الکترونیکی مورد تصریح قرار گرفته است. مواد مرقوم ذخیره، پردازش و توزیع داده های شخصی اشخاص را تنها در محدوده رضایت وی، اهداف مشخص، به صورت کاملاً شفاف به اندازه متناسب تعیین نموده و به مالکان اطلاعات این امکان را داده است تا در هر زمان امکان دسترسی به اطلاعات مذکور و حذف آنها از فرایند پردازش یا ذخیره و تبادل را داشته باشند. بدیهی است با عنایت به مقررات عام مسئولیت مدنی، توجهاً به امره بودن مقررات موضوع مواد ۵۸ و ۵۹ قانون تجارت الکترونیکی نقض هر یک از شرایط مواد مرقوم می تواند منجر به مسئولیت مدنی ناقضان نیز می گردد.

۲-۳- مسئولیت پذیری نهادهی:

چالش حقوقی دیگر در سازوکار عملکرد هوش مصنوعی مسئولیت پذیری نهادهای فعال در این فرایند در جبران خسارات وارده می باشد. آیین نامه و مصوب ۲۰۱۶ در این زمینه واجد مقرراتی است. مواد ۴۵ و ۲۴ از مقررات مذکور ارائه تضامین مالی برای جبران خسارات ناشی از نقض قواعد امنیتی و ارائه گزارش سازوکار پردازش اطلاعات از سوی کنترل کننده به کشور متبوع دارنده اطلاعات را پیش بینی نموده است. علاوه بر آن، پروتکل الحاقی به ماده ۴۵ مقررات مرقوم واجد شرایطی برای کشورهای غیر اروپایی میزبان اطلاعات خصوصی اروپاییان از جمله شرط «کفایت» می باشد. مطابق با بخش اول از بند دوم پروتکل مزبور، کشورهای میزبان اطلاعات خصوصی اروپاییان که خارج از اتحادیه اروپا می باشند، در دریافت این اطلاعات باید واجد سطح امنیتی لازم برای حفظ امنیت داده های مذکور و شفافیت کافی برای امکان نظارت بر عملکرد نهادهای فعال در صلاحیت سرزمینی خود باشند. این امر در صورتی محقق خواهد شد که معاهدات یا دیگر مقررات مصوب اتحادیه اروپا از جمله مفاد دستورالعمل حمایت از حقوق جمعی مصرف کنندگان مصوب ۲۰۱۸ این

اتحادیه در پیش مسئولیت تضامنی دولت در موارد نقض قواعد امنیتی پردازش داده‌های خصوصی توسط شرکت‌های تبعه این کشور، مورد پذیرش کشور ثالث قرار گیرد.

علاوه بر آن چالش دیگر ضرورت سیاستگذاری تقنینی در پیاده‌سازی این سازوکار در نظام داخلی کشورهای عضو اتحادیه است. به عبارتی در صورتی که کشوری از مقررات مواد ۲۴ و ۴۵ آیین‌نامه پیروی ننماید، ضمانت اجرای این امر در چه می‌باشد؟ آیا امکان محکوم نمودن دولت مذکور در دادگاه‌های اتحادیه اروپا با سازوکاری معین وجود دارد؟ در این خصوص به نظر نگارندگان با اخذ وحدت ملاک از ماده ۱۶ دستورالعمل ۲۰۱۸ که در راستای مقررات ماده ۸۰ آیین‌نامه ۲۰۱۶ در حوزه اقامه دعای جمعی برای جبران خسارات ناشی از نقض قواعد امنیتی تصویب شده است، امکان پیش‌بینی مسئولیت تضامنی دولت‌های عضو اتحادیه و سازمان‌های فعال در محدوده صلاحیت سرزمینی آنها در جبران خسارات وارده ناشی از نقض قواعد تعیین شده در مقررات مذکور موجود است، چراکه عدم اجرای مقررات اولیه مواد ۴۶ و ۴۵ از حیث قواعد عام حقوقی به منزله تقصیری تلقی می‌گردد که خسارات وارده را منتسب دولت متبوع سازمان واردکننده زیان می‌گرداند. از طرف دیگر مبانی تصویب مقررات ماده ۱۶ حمایت از حقوق دارندگان اطلاعات و تسری این مقررات به موارد مشابهی که ضرورت حفاظت از حقوق دارندگان اطلاعات است در اتحادیه اروپا احساس می‌گردد خالی از هرگونه ایراد خواهد بود.

مضاف بر آنچه بیان شد، در خصوص مفاد پروتکل الحاقی ماده ۴۵ نیز می‌توان بیان داشت، وجود چنین سازوکاری اگرچه در جهت حفظ امنیت اطلاعات اروپاییان می‌تواند مفید باشد، اما برای کشورهای خارج از اتحادیه اروپا واجد چالش‌هایی است. اولاً حاکمیت و استقلال هیچ کشوری نخواهد پذیرفت که الزامات امنیتی کشور یا اتحادیه دیگری در نظام قانونگذاری داخلی آن کشور وارد گردد، ضمن این‌که سؤال پیش رو این است که معیار تعیین سطح امنیتی لازم که در بند دوم از پروتکل الحاقی مورد تصریح قرار گرفته است چه می‌باشد؟ به جهت آنکه در هر حال امکان تفاسیر متعدد از کیفیت تعیین امنیت لازم برای حفاظت از داده‌ها وجود دارد، به نظر نمی‌رسد معیار مشخصی نیز پیش روی اتحادیه اروپا در همکاری مبادلاتی با کشورهای دیگر جهان وجود داشته باشد. ثمره این امر ایجاد رویکردهای متعدد در مواجهه با این حکم خواهد بود؛ از طرف دیگر برخی از کشورهای در حال توسعه مانند ایران، حتی فاقد قوانین اولیه در زمینه حفاظت از داده‌های الکترونیکی مورد تبادل در مبادلات فرامرزی می‌باشند که پذیرش شرایط اتحادیه اروپا در نظام داخلی این کشورها جزء دیگر مشکلات اجرایی پیش روی اتحادیه خواهد بود (صادقی و ناصر، ۱۳۹۹: ۱۴-۱).

مسئله اصلی در این تحقیق، بررسی و تحلیل این الزامات اخلاقی و چگونگی تأثیر آنها بر فرآیند تصویب قانون هوش مصنوعی در اتحادیه اروپا است. این مسئله با توجه به پیچیدگی‌های فنی و اخلاقی هوش مصنوعی و نقش کلیدی آن در آینده جوامع، از اهمیت ویژه‌ای برخوردار است.

هدف از این تحقیق، ارائه دیدگاهی جامع و منسجم از چگونگی برخورد قانون‌گذاران اروپایی با این چالش‌ها و تأثیرات آنها بر سیاست‌گذاری‌های مرتبط با هوش مصنوعی است.

هوش مصنوعی با همه پتانسیل‌ها و فرصت‌هایی که به ارمغان می‌آورد، چالش‌های اخلاقی خود را دارد، بنابراین پرواضح است رعایت الزامات اخلاقی در طراحی و استفاده از این فناوری نه تنها به جلوگیری از پیامدهای منفی کمک می‌کند، بلکه باعث افزایش اعتماد و پذیرش عمومی نسبت به هوش مصنوعی می‌شود. لذا با توجه به سرعت پیشرفت تکنولوژی لازم

است که این الزامات اخلاقی همواره مورد بازنگری و به روزرسانی قرار گیرند تا همسو همگام با تحولات نوین بتوان بهترین استفاده از هوش مصنوعی را برد. اتحادیه اروپا یک استراتژی صریح برای هوش مصنوعی اتخاذ کرد و گروه تخصصی سطح بالایی را در زمینه هوش مصنوعی متشکل از ۵۲ عضو را برای ارائه مشاوره در مورد سرمایه‌گذاری و مسائل اخلاقی، حکومتی در رابطه با هوش مصنوعی در اروپا منصوب کرد. در آوریل ۲۰۱۹ گروه متخصصین دستورالعمل‌های اخلاقی برای هوش مصنوعی قابل اعتماد را منتشر کرد دستورالعمل‌ها به مسائل مربوط به مسئولیت، شفافیت و حفاظت از داده‌ها به عنوان بخش‌های کاملاً مرکزی توسعه هوش مصنوعی نگارش شده‌اند تا بتوان قابلیت اطمینان را در میان مخاطبان دوچندان سازد. (Cath et al., 2018: 510)

در چند سال گذشته، تعدادی دستورالعمل اخلاقی در رابطه با هوش مصنوعی توسعه یافته توسط شرکت‌ها، انجمن‌های تحقیقاتی و نمایندگان دولت تدوین گردید و بسیاری از آنها تا حدی با قوانین موجود هم پوشانی دارند، اما اغلب مشخص نیست که چگونه قوانین و دستورالعمل‌ها برای تعامل دقیق‌تر در نظر گرفته شده‌اند. به طور خاص، روشی که اصولاً در نظر گرفته شده است تا دیدگاه‌ها را اجرا کنند، اغلب نامشخص است. به عبارت دیگر، دستورالعمل‌های اخلاقی بر دیدگاه‌های هنجاری تمرکز دارند، اما اغلب از منظر رویه‌ای ضعیف هستند. دستورالعمل‌های اخلاقی گروه کارشناسی کمیسیون اتحادیه اروپا نشانه روشنی از چالش حاکمیتی مداوم برای اتحادیه اروپا و کشورهای عضو آن است. جالب اینجاست که اورزولا فن در لاین، (رئیس کمیسیون اتحادیه اروپا)، در طول کاندیداتوری خود اظهار داشت که در ۱۰۰ روز اول ریاست جمهوری خود، پیشنهاد‌های قانونی را برای رویکرد هماهنگ اروپایی به پیامدهای انسانی و اخلاقی هوش مصنوعی ارائه خواهد کرد. در نتیجه، در فوریه ۲۰۲۰ کمیسیون اروپا یک استراتژی دیجیتال شامل پیشنهادهایی برای توانمندسازی، تعالی و اعتماد به هوش مصنوعی و یک کتاب سفید در مورد هوش مصنوعی منتشر کرد. در عین حال، برداشت کمیسیون اتحادیه اروپا در مورد توسعه و مدیریت هوش مصنوعی نشان دهنده یک روند جهانی در رویکردهای دولتی و قضایی برای مشاهده منافع اجتماعی و صنعتی با هوش مصنوعی در کنار نگرانی‌های اخلاقی و قانونی است که باید مورد توجه و کنترل قرار می‌گرفت. در واقع هوش مصنوعی مفهومی از توسعه و حاکمیت با پتانسیل و یا ریسک خطر بالا بود زیرا از فناوری‌های نوظهور است و تجربه و حد و مرز و باید و نبایدهای خاص خود را می‌طلبد. بخشی از چالش اتحادیه اروپا مسلماً شامل ایجاد تعادل در مقررات در برابر اعتمادی است که در نوآوری فنی و توسعه اجتماعی وجود دارد، اما هوش مصنوعی به آن توجه می‌کند (Chatila & Havens, 2019: 15).

آنچه از تعریف ماشین در جهان ارائه گردیده است همان چیزی است که کار را آسان می‌کند، بنابراین مطلوب نیست که با مقررات نامتعادل یا عجولانه معرفی شده خطر تضعیف شود. البته باید در نظر داشت همانطور که استفاده اجتماعی و وابستگی به هوش مصنوعی و یادگیری ماشینی در حال افزایش است، جامعه به طور فزاینده‌ای نیاز به درک هر گونه پیامدها و خطرات منفی، نحوه توزیع منافع و قدرت و نیازهای حکمرانی قانونی را دارد (Doshi-Velez & Kim, 2017).

مسائل اخلاقی در هوش مصنوعی

اگر به بحث‌هایی درباره هوش مصنوعی و اخلاقیات که در عرصه جهانی ارایه می‌شود نگاهی بیندازیم، می‌توان دریابیم که در حال حاضر موضوعی پر جنب و جوش در میان دانشگاہیان و نهادهای سیاست‌مدار است. دستورالعمل‌های اخلاقی به ویژه، به عنوان یک ابزار حکمرانی، در چند سال گذشته شاهد توسعه بسیار قوی بوده‌اند.

به عنوان مثال، مطالعه‌ای درباره چشم‌انداز اخلاقی هوش مصنوعی جهانی که در سال ۲۰۱۹ منتشر شد، ۸۴ سند حاوی اصول یا دستورالعمل‌های اخلاقی برای هوش مصنوعی را شناسایی کرد. این مطالعه به این نتیجه رسید که حداقل در پنج رویکرد مبتنی بر اصول اخلاقی در سطح جهانی اتفاق نظر نسبی وجود دارد: شفافیت، عدالت و انصاف، استفاده غیر مضر، مسئولیت، و یکپارچگی و حفاظت از داده‌ها تمامی این موارد به اهمیت اخلاق و توجه ویژه اتحادیه و متخصصین به این مهم دارد در عین حال، مطالعات نشان می‌دهد که تفاوت‌های قابل توجهی در نحوه تفسیر این اصول وجود دارد. که چرا آنها مهم تلقی می‌شوند. رایج‌ترین اصل «شفافیت» است که به نظر می‌رسد مفهومی چندوجهی است. در همین حال، محقق اخلاق، هاگندورف معتقد است که نقطه ضعف دستورالعمل‌های اخلاقی این است که اخلاق هوش مصنوعی مکانیسم‌هایی برای ایجاد انطباق یا اجرای ادعاهای هنجاری خود ندارد به گفته هاگندورف، به همین دلیل است که اخلاق برای بسیاری از شرکت‌ها و مؤسسات جذاب است. وقتی شرکت‌ها و مؤسسات تحقیقاتی دستورالعمل‌های اخلاقی خود را تدوین می‌کنند، مکرراً ملاحظات اخلاقی را معرفی می‌کنند، یا تعهدات اخلاقی خود را اتخاذ می‌کنند، هاگندورف استدلال می‌کند که این امر با معرفی چارچوب‌های قانونی الزام‌آور واقعی مقابله می‌کند. بنابراین او تأکید زیادی بر اجتناب از مقررات به طور خاص به عنوان هدف اصلی دستورالعمل‌های اخلاقی صنعت هوش مصنوعی دارد (IEEE Standards Association, 2019).

اخلاق به عنوان ابزاری برای حکمرانی در توسعه هوش مصنوعی به شدت مورد تأکید قرار گرفته است. حتی اگر خودتنظیمی مطمئناً به عنوان استدلالی برای اجتناب از مداخله قوانین مشخص استفاده می‌شود، هنوز سؤال این است که آیا رشد سریع حوزه هوش مصنوعی به طور خاص به همان اندازه نقش مهمی در نتیجه‌گیری که این زمینه خاص نیاز دارد ایفا نمی‌کند. (Mehrabi et al., 2021: 35)

دستورالعمل‌های اخلاقی تأثیر روشنی بر کتاب سفید بعدی در مورد هوش مصنوعی از کمیسیون اتحادیه اروپا داشته است اما هنوز باید دید که همه این منابع چه نوع اهمیت و تأثیری بر توسعه هوش مصنوعی اروپا خواهند داشت. دستورالعمل‌های اخلاقی اشاره می‌کند که هوش مصنوعی قابل اعتماد دارای سه جزء است که باید در کل چرخه حیات هوش مصنوعی وجود داشته باشد:

- الف- باید قانونی و مطابق با کلیه قوانین و مقررات قابل اجرا باشد.
- ب- باید اخلاقی باشد و از رعایت اصول و ارزش‌های اخلاقی محافظت کند.
- ج- باید هم از نظر فنی و هم از نظر اجتماعی قوی باشد، زیرا سیستم‌های هوش مصنوعی علی‌رغم نیت خوب می‌توانند باعث آسیب غیرعمدی شوند.

یافته‌ها

رهنمودها بر مسائل اخلاقی و استحکام تمرکز داشتند، اما مسائل حقوقی را خارج از دستورالعمل‌های صریح قرار می‌دهند. همچنین گروه متخصص چهار اصل اخلاقی را ارائه می‌دهد که «بنیان» هوش مصنوعی قابل اعتماد را تشکیل می‌دهد: احترام به استقلال انسان، پیشگیری از آسیب، انصاف و توضیح‌پذیری

با این حال، برای تحقق رسیدن به هوش مصنوعی قابل اعتماد، آنها به هفت پیش‌نیاز اصلی پرداختند که به گفته آنها باید به‌طور مداوم در طول چرخه حیات سیستم هوش مصنوعی ارزیابی و مدیریت شوند:

عاملیت و نظارت انسانی- استحکام فنی و ایمنی- حریم خصوصی و حاکمیت داده- شفافیت- تنوع، عدم تبعیض و انصاف- رفاه اجتماعی و محیطی- مسئولیت پذیری

همانطور که در بالا ذکر شد، اگرچه دستورالعمل‌ها بر اخلاق و استحکام تأکید دارند و نه بر مسائل قانونی، اما جالب توجه است که عدم تبعیض و حفاظت از حریم خصوصی بعدها به عنوان دو مورد از هفت محور اصلی توسعه یافته‌تر و قانون توجه شدند. پیش نیازهای اخلاقی برای اجرای هوش مصنوعی قابل اعتماد در رابطه با توصیه‌های سرمایه‌گذاری و سیاست‌گذاری که توسط گروه متخصص نیز منتشر شده است ویژگی‌هایی مانند رویکرد مبتنی بر ریسک را توصیه می‌کند که هم متناسب و هم مؤثر در تضمین قانونی، اخلاقی و قوی بودن هوش مصنوعی در انطباق با حقوق اساسی است. جالب توجه است، این گروه متخصص خواستار انجام نقشه‌برداری جامع از مقررات مربوطه اتحادیه اروپا است تا میزانی که مقررات مختلف هنوز اهداف خود را در دنیای مبتنی بر هوش مصنوعی برآورده می‌کنند، ارزیابی کنند. آنها تأکید می‌کنند که اقدامات قانونی و مکانیسم‌های کنترلی جدید ممکن است برای محافظت از حفاظت کافی در برابر اثرات منفی، و امکان نظارت و اجرای صحیح مورد نیاز باشد.

کتاب سفید در مورد هوش مصنوعی

همانطور که گفته شد، رئیس اتحادیه اروپا در دستورالعمل‌های سیاسی خود اعلام کرده بود یک رویکرد هماهنگ اروپایی در مورد مفاهیم انسانی و اخلاقی هوش مصنوعی و همچنین تأملی در مورد استفاده بهتر از هوش مصنوعی مدنظر قرار دارد تمامی راهکارها، بررسی‌ها، خطرات احتمالی، آزمایشات، موفقیت‌ها و مصوبات اخلاقی و قانونی به وسیله کارشناسان در کتابی با عنوان کتاب سفید بیان شده است که کمیسیون از یک رویکرد نظارتی و سرمایه محور با آنچه که «هدف دوگانه ارتقاء جذب هوش مصنوعی و پرداختن به خطرات مرتبط با استفاده‌های خاص از این فناوری جدید» می‌خواند، حمایت می‌کند. هدف از کتاب سفید این است که گزینه‌های سیاستی در مورد چگونگی دستیابی به این اهداف را تعیین کند. یک پیشنهاد کلیدی در کتاب سفید اتخاذ رویکردی مبتنی بر ریسک و بخش خاص برای تنظیم هوش مصنوعی است که در آن برنامه‌های پرخطر از همه برنامه‌های کاربردی دیگر متمایز می‌شوند. به عنوان مثال توسط کمیسیون اخلاق داده آلمان در کتاب سفید یک مبنای ارزشی مشخص با تمرکز ویژه بر مفهوم اعتماد وجود دارد: «با توجه به تأثیر عمده‌ای که هوش مصنوعی می‌تواند بر جامعه ما داشته باشد و نیاز به ایجاد اعتماد، بسیار حیاتی است که هوش مصنوعی اروپا بر این اساس استوار باشد. ارزش‌ها و حقوق اساسی ما مانند کرامت انسانی و حفاظت از حریم خصوصی هدف مشخص از چارچوب سیاست اتحادیه اروپا بسیج منابع برای دستیابی به یک اکوسیستم برتری در امتداد "کل زنجیره ارزش" بود». امیدواری کمیسیون این است که یک چارچوب نظارتی روشن اروپایی اعتماد بین مصرف‌کنندگان و کسب‌وکارها را در هوش مصنوعی ایجاد کند و در نتیجه جذب این فناوری را تسریع بخشد. (Floridi et al., 2018: 690)

کتاب سفید آدرسی روشن به رویکرد انسان محوری بر اساس ارتباطات ایجاد اعتماد در هوش مصنوعی انسان محور دارد، که همچنین بخش مرکزی دستورالعمل‌های اخلاقی است که در بالا مورد بحث قرار گرفت. کتاب سفید بیان می‌کند که کمیسیون ورودی‌های به دست آمده در مرحله آزمایشی دستورالعمل‌های اخلاقی تهیه شده توسط گروه کارشناسی سطح بالا در زمینه هوش مصنوعی را در نظر خواهد گرفت. قید این موضوع که در موضع فناوری و قانون و تعادل هارمونیک آن تاریخ به ما می‌آموزد که تعادل نظارتی دشوار است، به ویژه در زمان تغییرات سریع تکنولوژیک در جامعه در عین حال، محققین حقوقی مانند کارل رنر که قوانین مالکیت صنعتی شدن اروپای غربی را تحلیل کرد، می‌توان نتیجه گرفت که قانون

می‌تواند ارگانیسمی بسیار پویا و سازگار باشد. و می‌توان تصور کرد که بخش‌های مرکزی دستورالعمل‌های اخلاقی با حمایت از قوانین و مقررات اروپایی و ملی، با تمرکز بر اهمیت اعتماد (اکوسیستم) رسمی شود. تفسیر از قوانین موجود با توجه به قابلیت‌ها، امکانات و چالش‌های سیستم‌های هوش مصنوعی نیز یک نگرانی جدی مرتبط با چالش‌های اصلی است. اکنون که قانون هوش مصنوعی اروپا مصوب و لازم‌الاجرا شده است اطلاع‌رسانی و دقت نظر و همچنین اظهارنظرهای گوناگون سایر کشورهای اروپایی فرآیندهای قانونی با هدف یادگیری فناوری‌ها با تعامل و مستمر خواهد بود و اینگونه نیست که نظارت کارآمد و قانون را تمام شده در نظر بگیرند. (Kilian, 2020: 22)

رعایت عدم تبعیض و انصاف

الگوریتم‌های هوش مصنوعی می‌توانند ناآگاهانه تعصبات و نابرابری‌های موجود در داده‌ها را تقویت کنند. به عنوان مثال، سیستم‌های تشخیص چهره ممکن است در شناسایی افراد با رنگ پوست تیره‌تر دقت کمتری داشته باشند، که منجر به تبعیض ناعادلانه می‌شود. (Mittelstadt et al., 2016)

در واقع مهمترین مسائل و چالش برانگیزترین موارد کارگروه تخصصی کمیسیون هوش مصنوعی اروپا به صورت اجمالی این موارد بودند:

حریم خصوصی و امنیت داده‌ها: هوش مصنوعی برای عملکرد خود به حجم بزرگی از داده‌ها نیاز دارد، که این امر می‌تواند به نقض حریم خصوصی افراد منجر شود. همچنین، نگهداری و استفاده از این داده‌ها نیازمند تدابیر امنیتی قوی است تا از سوءاستفاده‌های احتمالی جلوگیری شود. با رشد سریع فناوری هوش مصنوعی^۱ مسائل مربوط به حریم خصوصی و امنیت داده‌ها به یکی از نگرانی‌های اصلی تبدیل شده است. هوش مصنوعی قادر به پردازش حجم عظیمی از داده‌ها و استخراج اطلاعات با ارزشی از آن‌هاست، اما همین قابلیت‌ها می‌توانند به تهدیدی برای حریم خصوصی و امنیت داده‌ها تبدیل شوند (Mohassel & Zhang, 2017: 20).

- جمع‌آوری گسترده داده‌ها: سیستم‌های هوش مصنوعی برای یادگیری و بهبود عملکرد خود نیاز به داده‌های زیادی دارند. این داده‌ها می‌توانند شامل اطلاعات حساس و شخصی کاربران باشند که در صورت سوءاستفاده، می‌توانند حریم خصوصی افراد را به خطر بیندازند.
- ناشناس‌سازی ناکافی: بسیاری از الگوریتم‌های هوش مصنوعی به داده‌های ناشناس نیاز دارند، اما فرآیند ناشناس‌سازی همیشه کارآمد نیست. داده‌های ناشناس‌شده می‌توانند از طریق ترکیب با سایر داده‌ها دوباره شناسایی شوند.
- پروفایلینگ و پیش‌بینی رفتار: الگوریتم‌های هوش مصنوعی قادر به ایجاد پروفایل‌های دقیق از افراد و پیش‌بینی رفتار آن‌ها هستند. این امر می‌تواند منجر به نقض حریم خصوصی و استفاده نادرست از اطلاعات برای اهداف تبلیغاتی یا نظارتی شود.
- حملات سایبری: سیستم‌های هوش مصنوعی می‌توانند هدف حملات سایبری قرار گیرند. هکرها می‌توانند با دسترسی به این سیستم‌ها، داده‌های حساس را به سرقت ببرند یا الگوریتم‌ها را دستکاری کنند.
- حملات تقلبی:^۲ در این نوع حملات، مهاجمان داده‌های ورودی را به گونه‌ای تغییر می‌دهند که سیستم هوش مصنوعی دچار خطا شود. این نوع حملات می‌تواند امنیت سیستم‌ها را به شدت تضعیف کند.

¹ AI

² Adversarial Attacks

- نقص‌های امنیتی در الگوریتم‌ها: الگوریتم‌های هوش مصنوعی ممکن است دارای نقص‌های امنیتی باشند که می‌توانند توسط مهاجمان مورد سوءاستفاده قرار گیرند. به عنوان مثال، نقص در الگوریتم‌های یادگیری ماشین می‌تواند منجر به نتایج نادرست یا غیرقابل اعتماد شود.
- استفاده از تکنیک‌های پیشرفته ناشناس‌سازی: به کارگیری تکنیک‌های پیشرفته ناشناس‌سازی داده‌ها می‌تواند به کاهش خطر شناسایی مجدد کمک کند. این تکنیک‌ها باید به گونه‌ای باشند که تعادل مناسبی بین حفظ حریم خصوصی و دقت الگوریتم‌های هوش مصنوعی برقرار کنند.
- پیاده‌سازی امنیت در سطح داده‌ها و الگوریتم‌ها: اتخاذ رویکردهای امنیتی چندلایه، از جمله رمزنگاری داده‌ها، احراز هویت قوی و نظارت مستمر بر الگوریتم‌ها، می‌تواند امنیت سیستم‌های هوش مصنوعی را افزایش دهد.
- آموزش و آگاهی‌بخشی: آموزش کاربران و توسعه‌دهندگان در زمینه حریم خصوصی و امنیت داده‌ها می‌تواند به کاهش خطرات کمک کند. آگاهی‌بخشی در مورد تهدیدات و راهکارهای موجود می‌تواند به تقویت فرهنگ امنیتی در سازمان‌ها منجر شود.
- رعایت اصول اخلاقی: تدوین و رعایت اصول اخلاقی در توسعه و استفاده از هوش مصنوعی می‌تواند به حفظ حریم خصوصی و امنیت داده‌ها کمک کند. این اصول باید شامل شفافیت، مسئولیت‌پذیری و احترام به حقوق کاربران باشند.
- حریم خصوصی و امنیت داده‌ها: در هوش مصنوعی موضوعی حیاتی و پیچیده است که نیازمند توجه جدی از سوی توسعه‌دهندگان، کاربران و قانون‌گذاران است. با پیاده‌سازی راهکارهای مناسب و رعایت اصول اخلاقی، می‌توان از مزایای هوش مصنوعی بهره‌مند شد و همزمان حریم خصوصی و امنیت داده‌ها را حفظ کرد. به همین منظور، همکاری بین‌المللی و تدوین قوانین و مقررات مناسب نیز ضروری به نظر می‌رسد.

۳- مسئولیت در هوش مصنوعی

سیستم‌های هوش مصنوعی باید شفاف و قابل توضیح باشند. مسئولیت نتایج و تصمیمات تولید شده توسط این سیستم‌ها باید مشخص باشد تا اطمینان حاصل شود که در صورت بروز خطا یا آسیب، افراد یا سازمان‌های مسئول پاسخگو خواهند بود. مسئولیت در هوش مصنوعی به معنای تعهد به پاسخگویی در قبال عملکرد، تصمیم‌گیری‌ها و تأثیرات سیستم‌های هوش مصنوعی است. مسئولیت و شفافیت در هوش مصنوعی از اهمیت بالایی برخوردارند و می‌توانند به ایجاد اعتماد و اطمینان در استفاده از این فناوری‌ها کمک کنند. با پیاده‌سازی راهکارهای مناسب و تعامل فعال با ذینفعان، می‌توان از مزایای هوش مصنوعی بهره‌مند شد و همزمان مسئولیت‌پذیری و شفافیت را تضمین کرد. این امر نه تنها به بهبود کیفیت و کارایی سیستم‌های هوش مصنوعی کمک می‌کند، بلکه از نگرانی‌های اجتماعی و اخلاقی نیز می‌کاهد این موضوع شامل جنبه‌های زیر می‌شود:

- مسئولیت اخلاقی: توسعه‌دهندگان و کاربران هوش مصنوعی باید در برابر تأثیرات اخلاقی و اجتماعی این فناوری‌ها پاسخگو باشند. این شامل اطمینان از عدم تبعیض، حفظ حقوق افراد و جلوگیری از سوءاستفاده‌های احتمالی است.
- مسئولیت قانونی: سیستم‌های هوش مصنوعی باید با قوانین و مقررات موجود سازگار باشند. این شامل رعایت قوانین حریم خصوصی، امنیت داده‌ها و مقررات مربوط به استفاده از فناوری‌های هوش مصنوعی در حوزه‌های مختلف می‌شود.

- مسئولیت فنی: توسعه‌دهندگان باید اطمینان حاصل کنند که سیستم‌های هوش مصنوعی از لحاظ فنی صحیح و قابل اعتماد هستند، این شامل تست‌های مداوم، بهبود الگوریتم‌ها و تضمین عملکرد صحیح سیستم‌ها در شرایط مختلف است (ISO/IEC JTC 1, 2020: 40).

۴- شفافیت در هوش مصنوعی

شفافیت به معنای قابلیت دسترسی و درک فرآیندها و تصمیم‌گیری‌های سیستم‌های هوش مصنوعی توسط کاربران و ذینفعان است. شفافیت می‌تواند به ایجاد اعتماد و کاهش نگرانی‌ها در مورد استفاده از هوش مصنوعی کمک کند. جنبه‌های مختلف شفافیت عبارتند از:

- توضیح‌پذیری: سیستم‌های هوش مصنوعی باید قادر به توضیح دلایل و فرآیندهای تصمیم‌گیری خود باشند. این امر می‌تواند به کاربران کمک کند تا تصمیمات این سیستم‌ها را درک کنند و به آن‌ها اعتماد بیشتری داشته باشند.
- دسترسی به داده‌ها و الگوریتم‌ها: شفافیت شامل ارائه دسترسی به داده‌های مورد استفاده و الگوریتم‌های به کار رفته در سیستم‌های هوش مصنوعی است. این امر می‌تواند به ارزیابی مستقل و نظارت بر عملکرد این سیستم‌ها کمک کند.
- گزارش‌دهی و مستندسازی: ارائه گزارش‌های جامع و دقیق از عملکرد سیستم‌های هوش مصنوعی و مستندسازی کامل فرآیندهای توسعه و استفاده از آن‌ها می‌تواند به افزایش شفافیت کمک کند.
- تدوین استانداردها و مقررات: ایجاد استانداردها و مقررات مشخص برای توسعه و استفاده از هوش مصنوعی می‌تواند به افزایش مسئولیت‌پذیری و شفافیت کمک کند. این استانداردها باید شامل مواردی مانند ارزیابی اثرات اجتماعی و اخلاقی، تضمین حریم خصوصی و امنیت داده‌ها و ارائه توضیحات کافی در مورد عملکرد سیستم‌ها باشند.
- آموزش و آگاهی‌بخشی: افزایش آگاهی و آموزش کاربران و توسعه‌دهندگان در مورد مسئولیت‌ها و نیازهای شفافیت در هوش مصنوعی می‌تواند به بهبود این موارد کمک کند. این شامل آموزش در مورد اصول اخلاقی، قوانین مرتبط و بهترین روش‌های توسعه و استفاده از هوش مصنوعی است.
- نظارت و ارزیابی مستقل: ایجاد نهادها و سازمان‌های مستقل برای نظارت و ارزیابی سیستم‌های هوش مصنوعی می‌تواند به افزایش شفافیت و مسئولیت‌پذیری کمک کند. این نهادها می‌توانند به بررسی عملکرد، ارزیابی تأثیرات اجتماعی و اخلاقی و ارائه توصیه‌های مناسب بپردازند.
- تعامل با ذینفعان: مشارکت و تعامل فعال با ذینفعان مختلف از جمله کاربران، سازمان‌های غیرانتفاعی، دولت‌ها و جوامع محلی می‌تواند به بهبود شفافیت و مسئولیت‌پذیری کمک کند. این تعامل می‌تواند به شناسایی نگرانی‌ها، دریافت بازخورد و توسعه راهکارهای مشترک منجر شود.

۵- قوانین و مقررات بین‌المللی

با پیشرفت سریع هوش مصنوعی و گسترش کاربردهای آن در زمینه‌های مختلف، نیاز به قوانین و مقررات بین‌المللی برای مدیریت توسعه و استفاده از این فناوری بیش از پیش احساس می‌شود. این قوانین و مقررات می‌توانند به تنظیم استانداردهای مشترک، تضمین امنیت و حریم خصوصی، و کاهش نابرابری‌ها و سوءاستفاده‌ها کمک کنند. در این مقاله به بررسی مهم‌ترین قوانین و مقررات بین‌المللی در حوزه هوش مصنوعی پرداخته می‌شود (مصطفوی و همکاران، ۱۴۰۲: ۹۰).

¹ Explainability

کشورهای مختلف در تلاش هستند تا چارچوب‌های قانونی مناسبی برای استفاده از هوش مصنوعی ایجاد کنند. اتحادیه اروپا با ارائه "راهنمای اخلاقی برای هوش مصنوعی" و ایالات متحده با توسعه قوانین مربوط به حفظ حریم خصوصی داده‌ها، نمونه‌هایی از این تلاش‌ها هستند از استانداردهای مورد توجه بین‌المللی می‌توان به موارد زیر اشاره کرد:

۱-۵- استانداردهای IEC، ISO:

سازمان بین‌المللی استانداردسازی^۱ و کمیسیون بین‌المللی الکتروتکنیک^۲ با همکاری هم در حال تدوین استانداردهای بین‌المللی برای هوش مصنوعی هستند. این استانداردها شامل اصول اخلاقی، امنیت، کیفیت داده‌ها و شفافیت الگوریتم‌ها می‌شوند.

راهنمای اخلاقی یونسکو: یونسکو در سال ۲۰۲۱ یک راهنمای اخلاقی برای هوش مصنوعی منتشر کرد که شامل اصولی نظیر احترام به کرامت انسانی، عدم تبعیض، شفافیت و پاسخگویی است. این راهنما به کشورهای عضو کمک می‌کند تا چارچوب‌های ملی خود را بر اساس این اصول تنظیم کنند.

۲-۵- چارچوب اخلاقی OECD:

سازمان همکاری و توسعه اقتصادی^۳ در سال ۲۰۱۹ اصولی برای هوش مصنوعی تدوین کرد که شامل شفافیت، پاسخگویی، امنیت، حریم خصوصی و تضمین حقوق انسانی می‌شود. این اصول به‌عنوان مرجعی برای کشورها و سازمان‌ها در تدوین مقررات ملی و بین‌المللی مورد استفاده قرار می‌گیرند.

۳-۵- حقوق و مسئولیت‌ها

یکی از مسائل کلیدی در حوزه هوش مصنوعی، تعیین حقوق و مسئولیت‌های مرتبط با استفاده از این فناوری است. این شامل حقوق کاربران در زمینه حریم خصوصی، حقوق توسعه‌دهندگان در زمینه استفاده از داده‌ها و مسئولیت‌های قانونی در صورت بروز خطاها و خسارات می‌شود. حقوق و مسئولیت‌ها در حوزه هوش مصنوعی از اهمیت ویژه‌ای برخوردارند و نیازمند توجه جدی از سوی توسعه‌دهندگان، کاربران و نهادهای نظارتی هستند. با تدوین قوانین و مقررات مناسب و ایجاد یک چارچوب اخلاقی قوی، می‌توان از مزایای هوش مصنوعی بهره‌مند شد و همزمان از چالش‌ها و مخاطرات مرتبط با آن جلوگیری کرد. همکاری بین‌المللی و تعامل فعال بین ذینفعان مختلف نیز می‌تواند به بهبود وضعیت حقوق و مسئولیت‌ها در این حوزه کمک کند. از جمله موارد مهم می‌توان به موارد زیر اشاره کرد:

- حقوق کاربران و حریم خصوصی: کاربران حق دارند که داده‌های شخصی‌شان به صورت امن و محرمانه نگهداری شود. توسعه‌دهندگان هوش مصنوعی باید از استانداردهای بالای امنیتی برای محافظت از داده‌ها استفاده کنند.
- شفافیت: کاربران باید حق داشته باشند که بدانند چگونه از داده‌هایشان استفاده می‌شود و الگوریتم‌های هوش مصنوعی چگونه تصمیم‌گیری می‌کنند. این شفافیت می‌تواند به افزایش اعتماد عمومی کمک کند.
- کنترل و انتخاب: کاربران باید قادر باشند که انتخاب کنند آیا داده‌هایشان برای آموزش و استفاده در سیستم‌های هوش مصنوعی مورد استفاده قرار گیرد یا خیر. این شامل گزینه‌های انصراف از استفاده نیز می‌شود.

^۱ ISO

^۲ IEC

^۳ OECD

۵-۴- حقوق توسعه‌دهندگان و شرکت‌ها

- مالکیت معنوی: توسعه‌دهندگان و شرکت‌ها حق دارند که از نتایج تحقیقات و نوآوری‌های خود بهره‌مند شوند و مالکیت معنوی آن‌ها حفظ شود. (Buolamwini & Gebru, 2018: 80)
- حمایت از نوآوری: قوانین باید به گونه‌ای باشد که از نوآوری و تحقیق و توسعه در زمینه هوش مصنوعی حمایت کند و موانع بی‌مورد را بر سر راه توسعه‌دهندگان قرار ندهد.
- عدم تبعیض: الگوریتم‌های هوش مصنوعی نباید به گونه‌ای طراحی شوند که تبعیض‌آمیز باشند. این شامل تبعیض بر اساس جنسیت، نژاد، مذهب، و دیگر ویژگی‌های شخصی می‌شود.
- امنیت و ایمنی: جامعه حق دارد که از امنیت و ایمنی سیستم‌های هوش مصنوعی اطمینان حاصل کند. این شامل حفاظت از زیرساخت‌های حیاتی و جلوگیری از سوءاستفاده‌های احتمالی است.

نتیجه گیری

اخلاق به عنوان یکی از پایه های اصلی در تدوین قوانین هوش مصنوعی عمل می کند. بدون در نظر گرفتن اصول اخلاقی، قوانین ممکن است ناکارآمد، ناعادلانه، یا حتی مضر باشند. برای مثال، قوانینی که به حریم خصوصی و حفاظت از داده های شخصی می پردازند، مستقیماً از اصول اخلاقی مرتبط با احترام به حقوق فردی و حریم خصوصی ناشی می شوند. در اتحادیه اروپا، تدوین قانون هوش مصنوعی با تأکید بر اصول اخلاقی مانند شفافیت، انصاف، و پاسخگویی انجام شده است. در بسیاری از موارد، رعایت قوانین هوش مصنوعی مستلزم رعایت اصول اخلاقی است. برای مثال، سیستم های هوش مصنوعی که در حوزه های حساس مانند سلامت، عدالت کیفری یا امور مالی استفاده می شوند، باید با اصول اخلاقی مطابقت داشته باشند تا اطمینان حاصل شود که این سیستم ها به طور عادلانه و بدون تبعیض عمل می کنند. در غیر این صورت، حتی اگر قوانین رعایت شوند، اما اصول اخلاقی نادیده گرفته شوند، ممکن است نتایج ناعادلانه ای به وجود آید که به بی اعتمادی عمومی و مشکلات اجتماعی منجر شود. یکی از چالش های اصلی در رعایت قانون هوش مصنوعی، تضاد میان اصول اخلاقی مختلف است. برای مثال، در برخی موارد ممکن است بین اصل شفافیت و حفظ حریم خصوصی تضاد وجود داشته باشد. همچنین، تعیین مسئولیت ها در مواقعی که یک سیستم هوش مصنوعی تصمیمات نادرست می گیرد، می تواند چالش برانگیز باشد. این تضادها نشان می دهند که رعایت قوانین هوش مصنوعی نیازمند موازنه ای دقیق میان اصول اخلاقی مختلف است. هوش مصنوعی با توانایی های گسترده خود می تواند تأثیرات عمیقی بر جوامع بشری داشته باشد. اما برای بهره برداری صحیح و مسئولانه از این فناوری، نیاز به تدوین و اجرای چارچوب های اخلاقی و قانونی مناسب است. با همکاری بین المللی و استفاده از تجارب مختلف، می توان اطمینان حاصل کرد که هوش مصنوعی به نفع همه افراد جامعه مورد استفاده قرار می گیرد. دشواری تعریف هوش مصنوعی آن را به عنوان یکی از چالش های نظارتی ناشی از پیاده سازی و توسعه هوش مصنوعی مطرح کرده است در حالی که میدان ناهمگن هوش مصنوعی از لحاظ تاریخی و معاصر شرایط مطلوبی را برای تحقیق و توسعه فراهم می کند، و مبهم بودن مفهومی چالشی را برای مقررات و حاکمیت ایجاد می کند. شاید این وابستگی به داده های یادگیری ماشینی امروزی است. قانون هوش مصنوعی اروپا نمونه ای از یک گام مهم در جهت ایجاد چارچوبی قانونی و اخلاقی برای توسعه و استفاده از هوش مصنوعی است. این قانون با هدف حفاظت از حقوق و آزادی های اساسی، افزایش شفافیت و اعتماد عمومی و تسهیل نوآوری تدوین شده است. با وجود چالش ها و انتقادات موجود، اجرای این قانون می تواند به بهبود کیفیت و مسئولیت پذیری در زمینه هوش مصنوعی کمک کند و نقش مهمی در شکل دهی آینده این فناوری در اروپا و جهان ایفا کند. لازم و بایسته است که سازمان ملی هوش مصنوعی ترتیبی اتخاذ نماید تا در فراخوان های علمی و تخصصی از علاقه مندان و متخصصان هوش مصنوعی در راستای نگارش قانون و رعایت اخلاق حرفه ای در این تکنولوژی نوپا و اجرای هدفمند آن اقدامات لازم را اعمال نماید. کاملاً پرواضح است دیری نمی پیماید که کاربردهای هوش مصنوعی در امور قضایی و به طور ویژه، امور قضاوت و وکالت به گونه ای پیش خواهد رفت که تمامی منابع انسانی موجود در دایره علم حقوق موظف به آموزش پذیری خواهند بود.

۱. جی گانکل، دیوید. (۱۴۰۱). حقوق ربات‌ها، ترجمه رضا فرج پور. تهران: انتشارات مجد.
۲. رهبری، ابراهیم و علی شعبانپور. (۱۴۰۱). چالش‌های کاربرد هوش مصنوعی به‌عنوان قاضی در دادرسی‌های حقوقی. فصلنامه تحقیقات حقوقی، ۲۵ (ویژه‌نامه حقوق و فناوری)، ۴۴۴-۴۱۹.
۳. <https://doi.org/10.52547/jlr.2022.228967.2335>
۴. مصطفوی اردبیلی، سید محمد مهدی؛ تقی‌زاده انصاری، مصطفی و رحمتی‌فر، سمانه. (۱۴۰۲). تأثیر هوش مصنوعی بر نظام حقوق بشر بین‌الملل. حقوق فناوری‌های نوین، ۴(۸)، ۱۰۰-۸۵.
۵. <https://doi.org/10.22133/mtlj.2023.378057.1149>
۶. صادقی، حسین و مهدی ناصر. (۱۳۹۹). چالش‌های اخلاقی و حقوقی آیین‌نامه اتحادیه اروپا در سازوکارهای هوش مصنوعی در حوزه سلامت. مجله اخلاق زیستی، ۱۰(۳۵)، ۱-۱۴.
۷. <https://doi.org/10.22037/bioeth.v10i35.27500>
8. Buolamwini, J., & Gebru, T. (2018). "Gender shades: Intersectional accuracy disparities in commercial gender classification". In *Conference on fairness, accountability and transparency*. PMLR. 77–91.
9. <https://proceedings.mlr.press/v81/buolamwini18a/buolamwini18a.pdf>
10. Cath, C., Wachter, S., Mittelstadt, B., Taddeo, M., & Floridi, L. (2018). "Artificial intelligence and the 'good society': the US, EU, and UK approach". *Science and engineering ethics*, 24, 505–528.
11. <https://doi.org/10.1007/s11948-017-9901-7>
12. Chatila, R., & Havens, J. C. (2019). "The IEEE global initiative on ethics of autonomous and intelligent systems". *Robotics and well-being*, 11–16.
13. https://doi.org/10.1007/978-3-030-12524-0_2
14. Doshi-Velez, F., & Kim, B. (2017). "Towards a rigorous science of interpretable machine learning". *arXiv preprint arXiv:1702.08608*.
15. <https://doi.org/10.48550/arXiv.1702.08608>
16. Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... & Vayena, E. (2018). "AI4People—an ethical framework for a good AI society: opportunities, risks, principles, and recommendations". *Minds and machines*, 28, 689–707.
17. <https://doi.org/10.1007/s11023-018-9482-5>
18. Gasser, U., & Almeida, V. A. (2017). A layered model for AI governance. *IEEE Internet Computing*, 21(6), 58–62.
19. <https://doi.org/10.1109/MIC.2017.4180835>
20. IEEE Standards Association. (2019). "IEEE P7001 Transparency of Autonomous Systems."
21. <https://doi.org/10.1109/IEEESTD.2022.9726144>
22. ISO/IEC JTC 1. (2020). "Information technology - Artificial intelligence - Overview of ethical and societal concerns." ISO/IEC TR 24028.
23. <https://iso.org/standard/78507.html>
24. Kilian, G. R. O. S. S. (2020). "White Paper on Artificial Intelligence-A European approach to excellence and trust". *European Commission*
25. <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52020DC0065&from=EN>
26. Lipton, Z. C. (2018). "The mythos of model interpretability: In machine learning, the concept of interpretability is both important and slippery". *Queue*, 16(3), 31–57.
27. <https://doi.org/10.1145/3236386.3241340>
28. Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). "A survey on bias and fairness in machine learning". *ACM computing surveys (CSUR)*, 54(6), 1–35.
29. <https://doi.org/10.1145/3457607>
30. Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). "The ethics of algorithms: Mapping the debate". *Big Data & Society*, 3(2), 2053951716679679.
31. <https://doi.org/10.1177/2053951716679679>
32. Mohassel, P., & Zhang, Y. (2017). "Secureml: A system for scalable privacy-preserving machine learning". In *2017 IEEE symposium on security and privacy (SP)* IEEE. 19–38.
33. <https://doi.org/10.1109/SP.2017.12>