

A Comparative Analysis of the Role and Requirements of Human Oversight on Artificial Intelligence in EU Law and Iranian Laws



Nahid Parsa

PhD in Private Law, University of Mazandaran,
Mazandaran, Iran nahidparsa84@yahoo.com



Abstract

The AI Act 2024 introduced by the European Commission highlights human oversight as a key safeguard for AI systems, setting a comprehensive legal framework for monitoring AI operations within the European Union. This legislation establishes clear requirements for the supervision of AI systems, particularly those categorized as high-risk. In contrast, Iranian law currently lacks specific regulations addressing human oversight in AI applications. However, indirect references to oversight responsibilities can be found in legal provisions such as Articles 7 and 12 of the Civil Liability Law, Article 333 of the Civil Code, Articles 69 and 112 of the Maritime Law, and Article 15 of the Mandatory Insurance Law of 2016.

Journal of Research and
Development in Public Law

Iranian Law and Legal Research
Institute

Vol. 1 | No. 2 | Fall 2024 and
Winter 2025
(Original Article)

<https://jrpl.illrc.ac.ir>

DOI:

<https://doi.org/10.22034/jrpl.2025.721655>

This article explores the nature of human oversight obligations in both legal systems by addressing what should be supervised, when oversight is required, and who bears the responsibility for supervision. It further highlights the uncertainties and gaps in the AI Act 2024 and the potential risks of excessive reliance on providers for ensuring the safety infrastructure of high-risk AI systems.

Keywords: AI, human-centric, human oversight, EU AI Act, automated decision-making



نقش و الزامات نظارت انسانی بر هوش مصنوعی در قانون اتحادیه اروپا و قوانین ایران

دکتری حقوق خصوصی دانشگاه مازندران، مازندران، ایران
nahidparsa84@yahoo.com

ناهید پارسا



دوفصلنامه تحقیق و توسعه در حقوق عمومی
پژوهشکده حقوق و قانون ایران

دوره ۱ | شماره ۲ | پاییز و زمستان ۱۴۰۳
(مقاله پژوهشی)

<https://jrpl.illrc.ac.ir>

DOI:
<https://doi.org/10.22034/jrpl.2025.721655>

چکیده

در قانون هوش مصنوعی ۲۰۲۴ اروپا، نظارت انسانی همچون یک اقدام حفاظتی اصلی برای اپراتورهای هوش مصنوعی تأکید شده است. کمیسیون اروپا با پیشنهاد قانون جدید، الزام دقیق و جامع نظارت انسانی بر سیستم‌های هوش مصنوعی را مطرح کرده است. در ایران، قانون مدونی برای نظارت انسان بر هوش مصنوعی وجود ندارد؛ اما با استفاده از موادی مانند مواد ۷ و ۱۲ قانون مسئولیت مدنی، ماده ۳۳۳ قانون مدنی، مواد ۶۹ و ۱۱۲ قانون دریایی و ماده ۱۵ قانون بیمه اجباری ۱۳۹۵، می‌توان موضع قانون‌گذار را در این زمینه مشخص کرد. این مقاله به سؤالاتی مانند چه چیزی باید نظارت شود، چه زمانی نظارت لازم است و چه کسی باید نظارت را انجام دهد، پاسخ می‌دهد. همچنین، به ابهامات و شکاف‌های قانون جدید هوش مصنوعی و پیامدهای اعتماد بیش‌ازحد به ارائه‌دهندگان برای ایمن‌سازی زیرساخت نظارتی سیستم‌های هوش مصنوعی پرخطر اشاره می‌کند.

این نوشتار با بررسی کتب و مقالات موجود و تجزیه و تحلیل آن‌ها، ساختاربندی شده است و پس از تعمق و غور در منابع موجود به این نتیجه می‌رسد که نه ماده ۱۴ قانون جدید هوش مصنوعی اروپا و نه قوانین ایران، جزئیات زیادی در مورد آنچه که ناظر انسانی در هنگام اجرای نظارت باید در نظر بگیرد یا توجه خود را به آن معطوف کند، بیان نمی‌کند. همچنین، نشان داد که فضای زیادی برای ارائه‌دهندگان باقی خواهد ماند تا تعیین کنند، جزئیات مربوط به چه داده‌هایی را بایستی به ناظران انسانی ارائه دهند. نتیجه‌گیری‌های مشابهی در مورد «زمان» اعمال نظارت انجام شد. ارائه‌دهندگان باید اطمینان حاصل کنند که کاربران در هر زمان ممکن است، فرایندهای سیستم را قطع یا لغو کنند.

کلیدواژه‌ها: هوش مصنوعی؛ انسان‌محوری، نظارت انسانی، قانون هوش مصنوعی اتحادیه اروپا، تصمیم‌گیری خودکار

پژوهشگاه علوم انسانی و مطالعات فرهنگی
پرتال جامع علوم انسانی

مقدمه

استفاده روزافزون از هوش مصنوعی، موضوعی داغ برای علاقه‌مندان و بحث‌های عمومی و دانشگاهی است. یکی از معضلات اساسی مرتبط با این بحث و توسعه این است که چگونه می‌توان از فناوری‌های هوش مصنوعی، بدون خطر یا آسیب‌رساندن به جامعه و شهروندان، استفاده کرد. در تلاش‌های اروپایی برای توسعه هوش مصنوعی، یک موضع مردد وجود دارد، زیرا درصد نسبتاً بالایی از محدودیت‌ها برای استقرار هوش مصنوعی یا الزامات حفاظتی مرتبط در زمان استفاده از آن‌ها، توسط اتحادیه ایجاد شده است. یک نقطه کانونی این است که چگونه می‌توان از کنترل معنادار انسانی در مکان و زمانی که نیاز است، اطمینان حاصل کرد. در این زمینه، «نظارت انسانی» معیاری است که بر آن تأکید شده است. کمیسیون اروپا در پیش‌نویس قانون جدید هوش مصنوعی اتحادیه اروپا از آوریل ۲۰۲۱، پیشنهاد کرده است تا الزاماتی را برای سیستم‌های هوش مصنوعی به اصطلاح پرخطر تعیین کند تا به گونه‌ای طراحی و توسعه داده شوند که بتوانند، در طول استفاده از آن‌ها به طور مؤثر توسط اشخاص حقیقی نظارت شود. این موضوع که در ماده ۱۴ قانون جدید هوش مصنوعی اروپا آمده است، نشان‌دهنده جاه‌طلبی و اعتماد به شایستگی نظارت انسانی است. این تجلی از لزوم رعایت اصول حفاظتی برای سلامت و ایمنی و حقوق اساسی انسان در اعلامیه کمیسیون و قانون جدید هوش مصنوعی اروپا در مورد هوش مصنوعی «انسان‌محور» است: هوش مصنوعی باید ابزاری برای مردم و نیروی خیر در جامعه باشد که هدف نهایی افزایش رفاه انسان است؛ بنابراین قوانین هوش مصنوعی موجود که در چارچوب اتحادیه و به اشکال دیگر بر افراد اتحادیه تأثیر می‌گذارند، باید مبتنی بر نظارت انسان باشد، طوری که مردم بتوانند به استفاده از این فناوری در قالبی ایمن و مطابق با قانون، از جمله احترام به حقوق اساسی، اعتماد کنند، بنابراین قانون جدید هوش مصنوعی اروپا و الزامات نظارتی آن در قوانین اروپا تازگی دارد. اینکه

1 new EU Artificial Intelligence Act1 from April 2021(AIA)- Council of Europe Ad hoc Committee on Artificial Intelligence (CAHAI), Revised Zero Draft [framework] Convention on Artificial Intelligence, Human rights, Democracy and the Rule of Law, 6 January 2023, CAI (2023)01.(CAI)

چه چیزی می‌تواند (یا باید) همچون عناصر اصلی چنین رویکردی در نظر گرفته شود، موضوعی است که به بحث‌های آکادمیک و همچنین سیاسی نیاز دارد. درحالی‌که الزام به تضمین نظارت انسانی مستلزم آن است که انسان‌ها برخی از خطرات مرتبط با سیستم‌های هوش مصنوعی را شناسایی کنند. این در حالی است که ظرفیت و استعداد آن‌ها برای انجام این کار به نوع سیستم‌هایی که باید نظارت شوند و همچنین شفافیت این سیستم‌ها بستگی دارد. علاوه بر این، وظایف و شرایط کاری آن‌ها (ازجمله آموزش‌هایی که می‌میانند) نیز حائز اهمیت است. توجه به چنین عواملی برای مشخص کردن انتظارات انسان‌ها از نظر کنترل بر سیستم‌های هوش مصنوعی مهم است. این امر، همچنین به این معنی است که «نوع» نظارتی که تجویز یا انجام می‌شود، می‌تواند بر نوع خطرات یا آسیب‌هایی که انسان‌ها قادر به شناسایی و کاهش آن هستند، تأثیر بگذارد.

در قوانین ایران به‌طور مستقیم به لزوم نظارت انسان بر اشیا، حیوان یا انسان دیگری، اشاره نشده است، اما مواد ۶۹ و ۱۱۲ قانون دریایی ایران و مواد ۷ و ۱۲ قانون مسئولیت مدنی و ماده ۱۵ قانون بیمه اجباری ۱۳۹۵، به لزوم نظارت یک انسان بر انسان دیگر، به‌طور تلویحی اشاره کرده‌اند و به مناسبت این وظیفه برای فرد، مسئولیت نیابتی قائل شده‌اند. همچنین، در ماده ۳۳۳ قانون مدنی به لزوم نظارت انسان بر اشیا اشاره شده است. حال شایسته است که در این مقاله، به‌طور جزئی‌تر، این مواد مورد بررسی قرار گیرند و تجزیه و تحلیل شود که آیا این مواد می‌توانند، در زمینه نظارت انسانی بر هوش مصنوعی، کاربرد داشته باشند یا خیر؟

در ادبیات حقوقی ایران، در رابطه با ضمان در هوش مصنوعی (رجبی، ۱۳۹۸) و هوش مصنوعی و حکمرانی آینده (قاسمی، ۱۴۰۰) و هوش مصنوعی و قانون‌گذاری (خوئی، ۱۳۹۶) مقالاتی نوشته شده است، اما در رابطه با نظارت انسان بر سیستم‌های هوش مصنوعی، تاکنون نوشتاری تهیه نشده است؛ بنابراین از نکات بدیع این تحقیق، بررسی چگونگی نظارت بر هوش مصنوعی با نگاهی به قانون هوش مصنوعی ۲۰۲۴ اروپا است. این مقاله به شرحی که در ادامه می‌آید، خواهد بود؛ ابتدا محتوا و پیکربندی ماده ۱۴ قانون جدید هوش مصنوعی اروپا و ارتباط آن با برخی دیگر از مقررات مربوط به نظارت انسانی را

معرفی می‌کند، سپس به اینکه بر چه چیزی، چه زمانی و توسط چه کسی، بایستی نظارت صورت پذیرد، می‌پردازد.

۱- مقدمه‌ای کوتاه بر ماده ۱۴ قانون هوش مصنوعی ۲۰۲۴ اروپا

یکی از جنبه‌های اصلی الزامات نظارت انسانی در ماده ۱۴ قانون جدید هوش مصنوعی اروپا، که دامنه کاربردی آن را محدود می‌کند، این است که فقط در مورد سیستم‌های هوش مصنوعی به اصطلاح با «خطر بالا» (و نه برای همه سیستم‌های هوش مصنوعی) اعمال می‌شود. نسخه پیش‌نویس پارلمان اروپا از قانون جدید هوش مصنوعی اروپا، گنجانیدن ماده ۴(الف) جدید را پیشنهاد می‌کرد که شامل یک اصل کلی از نظارت انسانی است و برای همه سیستم‌های هوش مصنوعی قابل اعمال باشد. با این حال، این امر هیچ‌گونه الزام نظارتی سختگیرانه‌ای را تحمیل نمی‌کند و فقط اپراتورها را مجبور می‌کند تا بیشترین تلاش خود را در توسعه و استفاده از سیستم‌های هوش مصنوعی یا مدل‌های بنیادی در راستای اصل «نماینده‌گی و نظارت انسانی» و سایر موارد انجام دهند؛ بنابراین، قانون ۲۰۲۴، این واقعیت را تغییر نمی‌دهد که رژیم نظارتی سختگیرانه قانون جدید هوش مصنوعی اروپا برای آن دسته از سیستم‌های هوش مصنوعی که با عنوان پرخطر طبقه‌بندی می‌شوند، محفوظ بماند. قانون هوش مصنوعی، میان آن دسته از سیستم‌های هوش مصنوعی که بالقوه خطرناک هستند (و این موضوع برایشان ممنوعیت همراه دارد) و سایر دسته‌ها تمایز قائل می‌شود. تقسیم‌بندی قانون هوش مصنوعی، شامل سیستم‌هایی که ریسک بالا، کم‌خطر یا حداقل خطر دارند، است. هیچ شمارش جامعی در قانون مذکور انجام نشده است، اما طبقه‌بندی «خطر بالا» در پیوست شماره ۳ این قانون آمده است. نمونه‌هایی از این سیستم‌ها از جمله مواردی که در ادامه می‌آیند، هستند؛ سیستم‌های هوش مصنوعی که برای شناسایی بیومتریک و دسته‌بندی افراد حقیقی، مدیریت و بهره‌برداری زیرساخت‌های حیاتی؛ آموزش و پرورش، اشتغال، مدیریت پرسنل و دسترسی به خوداشتغالی؛ دسترسی و استفاده از خدمات و مزایای اولیه خصوصی و عمومی؛ اجرای قانون؛ مهاجرت، پناهندگی و کنترل مرزها و اجرای عدالت و فرایندهای دموکراتیک استفاده می‌شوند، از جمله این دسته‌بندی‌ها به‌شمار می‌روند. با این حال، سخت‌گیری

طراحی قانون جدید هوش مصنوعی اروپا در این زمینه با این واقعیت که ماده ۷ قانون جدید هوش مصنوعی اروپا به کمیسیون این اختیار را می‌دهد تا به‌روزرسانی لیست سیستم‌های «پرخطر» در پیوست ۳ تا حدودی کاهش دهد؛ بنابراین با نمایندگی کمیسیون، انواع جدید یا نادیده گرفته شده از هوش مصنوعی که می‌توانند با عنوان پرخطر طبقه‌بندی شوند، در صورت شناسایی، به لیست اضافه خواهند شد. شمارش نسبتاً راکد انواع سیستم‌هایی که در دسته‌بندی «پرخطر» در نظر گرفته می‌شوند، به این معنی است که کاربرد الزامات نظارت انسانی در ماده ۱۴ قانون جدید هوش مصنوعی اروپا به پیچیدگی خاص سیستمی که در حال استقرار است بستگی ندارد؛ چالش‌های انجام نظارت انسانی ممکن است در رابطه با انواع سیستم‌ها و کاربردهای آن‌ها، متفاوت باشد (Enarsson, 2022: 123).

باتوجه به ساختار و محتوای ماده ۱۴ قانون جدید هوش مصنوعی اروپا، این ماده شامل پنج بند است که چهار بند اول مربوط به همه سیستم‌های هوش مصنوعی پرخطر است. پاراگراف اول - ۱۴(۱) - که قاعده اصلی از پیش معرفی شده را تعیین می‌کند، بیان می‌دارد: «سیستم‌های هوش مصنوعی پرخطر باید به گونه‌ای با ابزارهای رابط انسان و ماشین، طراحی و توسعه داده شوند که در طول دوره‌ای که سیستم در حال استفاده است، بتوان به‌طور مؤثر توسط اشخاص حقیقی بر آن‌ها نظارت کرد. همان‌طور که برخی اشاره کرده‌اند (Enqvist, 2023: 510-512)، این تعهد در بند ۱۴(۲) بیان شده است و مقرر می‌دارد، هدف نظارت باید پیشگیری یا به حداقل رساندن آن خطرات سلامت، ایمنی یا حقوق اساسی باشد که ممکن است هنگام استفاده از سیستم‌های هوش مصنوعی پرخطر مطابق با هدف موردنظر آن‌ها یا تحت شرایط سوءاستفاده، قابل پیش‌بینی باشد. این امر، به‌ویژه زمانی صدق می‌کند که این خطرات علی‌رغم اعمال سایر تعهدات حفاظتی اصلی که بر عهده ارائه‌دهندگان سیستم‌های هوش مصنوعی پرخطر گذاشته شده است - در همان فصل (۲) پیش‌نویس ارائه شده‌اند - رخ می‌دهند. در ماده ۱۰ اطمینان حاصل می‌شود که شیوه‌های مدیریت و حاکمیت داده‌ها برای آموزش، اعتبارسنجی و آزمایش مدل‌ها و مجموعه‌های داده وجود داشته باشد. در ماده ۱۱ اطمینان حاصل می‌شود که

اسناد فنی کافی در دسترس است. در ماده ۱۲ برای اطمینان از نگهداری صحیح سوابق؛ در ماده ۱۳ برای اطمینان از طراحی سیستم شفاف و دستورالعمل‌های مناسب برای استفاده؛ و در ماده ۱۵ برای اطمینان از صحت، استحکام و امنیت سایبری سیستم‌ها، قواعدی تعیین شده است. جنبه‌های طراحی فوق‌الذکر ماده ۱۴ قانون جدید هوش مصنوعی اروپا نشان‌دهنده یک ارتباط متقابل نظارتی و همچنین عملکردی با سایر اقدامات حفاظتی اصلی در مواد ۱۰-۱۲ و ۱۴-۱۵ قانون جدید هوش مصنوعی اروپا است. تعهد ارائه‌دهنده برای اطمینان از قابلیت‌های نظارت افزایش می‌یابد، به‌خصوص در جایی که نمی‌توان انتظار داشت، عملکردهای کاهش خطر سایر اقدامات اصلی حفاظتی به‌اندازه کافی کارآمد باشد. از نظر عملکردی، همه این اقدامات حفاظتی از اهداف شفافیت عمومی مقررات حمایت می‌کنند و همگی به اجزای مربوطه به زیرساخت‌های نظارتی در دسترس ناظر انسانی کمک خواهند کرد و از طریق واسطه‌های انسانی در قامت مفسر داده‌های تولیدشده توسط الگوریتم‌های عملیاتی، نظارت انسانی انجام می‌شود. بدون سطح کافی از شفافیت سیستم (همان‌طور که توسط سایر اقدامات حفاظتی حامی شفافیت، مانند مستندات، نگهداری سوابق و ارائه اطلاعات به کاربران و... ثابت شده است)، ناظران انسانی هیچ‌چیز اساسی برای بررسی ندارند. (Green, 2022:12) این روابط متقابل حقوقی و همچنین کارکردی میان ماده ۱۴ و سایر اقدامات حفاظتی قانون جدید هوش مصنوعی اروپا تأکید می‌کند که محتوا و پیامدهای این ماده بدون در نظر گرفتن برخی از سایر مقررات آیین‌نامه به‌طور کامل قابل تشریح یا درک نیست. چنین ملاحظاتِ همچنین ممکن است، اطلاعات بیشتری در مورد «چه چیزی»، «چه زمانی» و «توسط چه کسی» ماده ۱۴ قانون جدید هوش مصنوعی اروپا، با جزئیات بیشتر ارائه دهد. اقدامات ملموس‌تری که از طریق آن‌ها، تعهد نظارت باید توسط ارائه‌دهندگان سیستم‌های هوش مصنوعی انجام شود، در ماده ۱۴ (۳) قانون جدید هوش مصنوعی اروپا فهرست شده است. که دو گزینه را مطرح می‌کند. یک گزینه برای ارائه‌دهندگان این است که از طریق شناسایی و ایجاد اقدامات نظارتی بر روی سیستم هوش مصنوعی پرخطر قبل از عرضه در بازار یا در خدمت (در صورت امکان) نظارت انسانی را تضمین کنند. از طرف دیگر، ارائه

دهندگان همچنین می‌توانند اقدامات مناسبی را که قرار است توسط کاربر اجرا شود، شناسایی کنند. قابل ذکر است، این ماده این تعهدات را به جای کاربران سیستم‌های هوش مصنوعی پرخطر بر عهده ارائه دهندگان می‌گذارد. «ارائه‌دهنده» به عنوان یک شخص حقیقی یا حقوقی، مقام عمومی، آژانس یا نهاد دیگری تعریف می‌شود که یک سیستم هوش مصنوعی را توسعه می‌دهد یا یک سیستم هوش مصنوعی را با هدف عرضه آن در بازار یا به خدمت گرفتن آن تحت نظر خود، توسعه داده است. در عوض، «کاربران» به عنوان هر شخص حقیقی یا حقوقی، مقام عمومی، آژانس یا نهاد دیگری تعریف می‌شود که از یک سیستم هوش مصنوعی تحت اختیار خود استفاده می‌کند. این واقعیت که تعهدات مندرج در ماده ۱۴ قانون جدید هوش مصنوعی اروپا بر عهده ارائه دهندگان گذاشته شده است، ماهیت عمدتاً پیشگیرانه این ماده را برجسته می‌کند. ارزیابی ریسک باید در مرحله طراحی سیستم انجام شده و از طریق نصب «قابلیت‌های» نظارت، قبل از ارائه سیستم به کاربران، مورد توجه قرار گیرد. براساس برخی نظرات (Smuha, 2021: 57) Binns, 2022: 197، درک کمیسیون از مفهوم "نظارت انسانی" به ویژه بر تفسیر، تعقیب یا اصلاح خروجی سیستم‌ها توسط انسان متمرکز است و این نظارت به مفاهیم گسترده‌تر مانند "نظارت سازمانی" تعمیم نمی‌یابد. به عبارت دیگر، این نظارت انسانی تنها به عملکرد فردی انسان‌ها در تفسیر و اصلاح خروجی‌های سیستم‌ها محدود است و شامل ساختارها و فرایندهای کلی سازمانی نمی‌شود. (Harasimiuk, 2021: 49) ویژگی‌های فنی خاص‌تری که ارائه‌دهندگان باید سیستم‌های هوش مصنوعی پرخطر خود را به آن‌ها مجهز کنند، در ماده ۱۴ (۴) قانون جدید هوش مصنوعی اروپا فهرست شده‌اند. این ماده ارائه‌دهندگان را موظف می‌کند تا اطمینان حاصل کنند که سیستم هوش مصنوعی پرخطر به افرادی که نظارت انسانی به آنها محول شده است، برحسب شرایط، این امکان را می‌دهد که: (الف) ظرفیت‌ها و محدودیت‌های سیستم هوش مصنوعی پرخطر را به طور کامل درک کند و بتواند عملکرد آن را به درستی نظارت کند، به طوری که علائم ناهنجاری، اختلال در عملکرد و عملکرد غیرمنتظره را بتوان در اسرع وقت شناسایی و برطرف کرد. (ب) اپراتورها از پیامدهای تکیه بیش از حد به خروجی تولید شده توسط یک سیستم هوش مصنوعی

پرخاطر ("سوگیری خودکار")، باخبر هستند به ویژه برای سیستم‌های هوش مصنوعی پرخاطر که برای ارائه اطلاعات یا توصیه‌ها استفاده می‌شوند؛ ج) قادر به تفسیر صحیح خروجی سیستم هوش مصنوعی پرخاطر با در نظر گرفتن ویژگی‌های سیستم و ابزارها و روش‌های تفسیری موجود باشد. (د) بتواند در هر موقعیت خاص تصمیم بگیرد که از سیستم هوش مصنوعی پرخاطر استفاده نکند یا خروجی سیستم هوش مصنوعی پرخاطر را نادیده بگیرد، یا معکوس کند. (ه) قادر به مداخله در عملکرد سیستم هوش مصنوعی پرخاطر یا قطع کردن سیستم از طریق دکمه "توقف" یا روشی مشابه باشد.

قطع‌نامه‌های سازمان ملل متحد و مقررات شورای عالی فضای مجازی هر دو نقش مهمی در تنظیم مقررات میان‌المللی و ملی در حوزه‌های فناوری اطلاعات و ارتباطات دارند، به‌ویژه در زمینه‌های مرتبط با هوش مصنوعی و حریم خصوصی. اشاره به این موارد می‌تواند زمینه‌ساز توسعه قواعد و استانداردهای دقیق‌تری برای نظارت انسانی بر سیستم‌های هوش مصنوعی باشد.

قطع‌نامه‌های سازمان ملل متحد

سازمان ملل متحد در زمینه فناوری‌های نوین و هوش مصنوعی چندین قطع‌نامه و توصیه‌نامه به تصویب رسانده است که هدف آن‌ها تضمین استفاده امن، اخلاقی و عادلانه از این فناوری‌ها در سطح جهانی است. مهم‌ترین موارد شامل:

۱ قطع‌نامه ۲۸۲/۷۲: (۲۰۱۸) در این قطع‌نامه، موضوع استفاده مسئولانه از هوش مصنوعی و فناوری‌های نوین در راستای حفاظت از حقوق بشر و رفاه اجتماعی مطرح شده است. به‌ویژه در این قطع‌نامه، تأکید شده که کشورهای عضو باید اخلاقی‌بودن و امنیت را در استفاده از هوش مصنوعی در نظر داشته باشند.

• قطع‌نامه ۲۱۹/۷۳: (۲۰۱۹) این قطع‌نامه به‌طور خاص بر هوش مصنوعی و حقوق بشر تمرکز دارد و از کشورهای عضو خواسته است که چارچوب‌های نظارتی ملی خود را در راستای حفاظت از حریم خصوصی و امنیت سایبری با توجه به فناوری‌های نوین، به‌ویژه هوش مصنوعی، تنظیم کنند.

• **قطع نامه ۱۴۳/۷۴: (2020)** این قطع نامه به ایجاد قوانین بین المللی برای مدیریت هوش مصنوعی و استفاده از آن در عرصه های مختلف از جمله بهداشت، آموزش و حمل و نقل می پردازد و بر نظارت انسانی و مسئولیت پذیری در این زمینه تأکید می کند. این قطع نامه ها به طور غیرمستقیم بر نظارت انسانی تأکید دارند، زیرا برای ایجاد تعادل میان نوآوری و امنیت، نظارت انسانی بر سیستم های هوش مصنوعی ضروری است.

۲- موضوع مورد نظارت توسط ناظران انسانی

در حقوق ایران، وظیفه نظارت انسانی بر اشیا در ماده ۳۳۳ قانون مدنی تعیین شده است که موضوع نظارت را دیوار و اشیا در نظر گرفته است. ماده ۱۵ قانون بیمه اجباری ۱۳۹۵ و ماده ۱۲ و ۷ قانون مسئولیت مدنی و ماده ۶۹ و ۱۱۲ قانون دریایی، همگی تعیین می کنند که انسان باید مورد نظارت قرار گیرد. ماده ۱۵ قانون بیمه اجباری در جواب آنچه که باید مورد نظارت قرار گیرد، اعمال کارآموز در طول دوران آموزش را در نظر می گیرد و ماده ۷ و ۱۲ قانون مسئولیت مدنی اعمال صغیر یا مجنون و کارگر را موضوع نظارت معرفی می کند. مواد ۶۹ و ۱۱۲ نیز اعمال کارکنان را موضوع نظارت می داند، اما در قانون جدید اتحادیه اروپا چه چیزی باید تحت نظارت و کنترل انسانی قرار گیرد؟

اهداف نظارت انسانی، تعامل میان ناظران انسانی و اطلاعات خاصی است که سیستم هوش مصنوعی به آن ها ارائه می دهد (Onitui, 2022: 175). ناظران انسانی برای اینکه بتوانند، خطرات، سوگیری ها یا خطاهای احتمالی را شناسایی کرده و به آن ها واکنش نشان دهند، نه تنها باید خروجی سیستم را در نظر بگیرند، بلکه باید آن را ارزیابی و تحلیل کنند. برای تعیین اینکه «چه چیزی» باید مورد نظارت قرار گیرد، جنبه های عملیاتی یک سیستم هوش مصنوعی پرخطر که قرار است توسط انسان نظارت و بازبینی شود، حائز اهمیت است. یک نقطه شروع مناسب برای پرداختن به «چه چیزی» ماده ۱۴ قانون جدید هوش مصنوعی اروپا، در نظر گرفتن اهداف اعلام شده آن است. به بیان دقیق، شاید بهتر است، اینها را در قامت سؤالاتی در مورد «چرا» نظارت انسانی لازم است، تلقی کرد. با این حال، هدف کلی این ماده برای جلوگیری یا به حداقل رساندن خطرات برای سلامت، ایمنی یا حقوق اساسی نیز نشان می دهد که نظارت باید در جهت شناسایی برخی اثرات

نامطلوب سیستم‌های هوش مصنوعی پرخطر بر منافع و ارزش‌های عمومی مهم باشد. البته ازمنظر نظارتی، توصیفات بسیار گسترده و جامع از آنچه ناظران انسانی در بررسی خود باید در نظر بگیرند، اهداف خاص‌تر نظارت را تحت الشعاع قرار می‌دهد. ماده ۱۴ قانون جدید هوش مصنوعی اروپا در ارتباط با سایر تعهدات حفاظتی اصلی که بر عهده ارائه‌دهنده گذاشته شده است، به‌طور غیرمستقیم بر اهمیت شفافیت سیستم - در قامت یک ضرورت برای توانمندسازی نظارت - تأکید می‌کند. این مقاله در درجه اول بر جنبه رابطه‌ای شفاف میان سیستم‌ها و انسان‌ها تأکید دارد؛ بنابراین شفافیت نه‌تنها همچون شکلی از «بازبودن» تعبیر می‌شود، بلکه در قامت کیفیت قابل‌شناسایی و قابل‌درک‌بودن نیز تعبیر می‌شود. این امر از طریق تعهدات ارائه‌دهنده برای قادرساختن ناظران انسانی به درک کامل ظرفیت‌ها و محدودیت‌های سیستم هوش مصنوعی پرخطر و نظارت مناسب بر عملکرد آن مشخص می‌شود؛ برای مثال درحالی‌که جزئیاتی در مورد اینکه چگونه خروجی یک سیستم را باید ارائه کرد تا اطمینان حاصل شود که انسان‌ها قادر به تفسیر صحیح آن هستند، معیاری ارائه نمی‌شود. این فرمول‌ها فقط به جنبه مرکزی هدف نظارت انسانی اشاره می‌کنند. اگر اطلاعات (یا خروجی سیستم) در چنین نمایشی به ناظر ارائه نشود که بتواند آن را تفسیر یا درک کند، انجام نظارت هیچ فایده‌ای ندارد و ازآنجایی‌که سیستم‌های هوش مصنوعی پرخطر از نظر پیکربندی، هدف و کاربرد، بسیار متفاوت خواهند بود، ممکن است ارائه‌دهندگان سیستم برای ارزیابی‌های مستدل از نوع اطلاعات و ارائه‌ای که به انتقال کیفی دانش به ناظران انسانی کمک می‌کند مجهزتر باشند. بااین‌حال، قانون جدید اروپا برای ارائه‌دهندگان فضایی اختیاری باقی می‌گذارد تا نحوه پیکربندی سیستم‌های هوش مصنوعی پرخطر خود را برای رسیدن به این اهداف تعیین کنند. همچنین، ماده ۱۴ قانون جدید هوش مصنوعی اروپا شامل تعهداتی است که به‌طور خاص به این موضوع می‌پردازد که چه نوع اختیاراتی را ناظران انسانی در رابطه با سیستم دارند. ارائه‌دهندگان سیستم‌های هوش مصنوعی باید مطمئن شوند که انسان‌ها می‌توانند در هر موقعیت خاص تصمیم بگیرند که از سیستم هوش مصنوعی پرخطر استفاده نکنند یا خروجی آن را نادیده بگیرند. آن‌ها همچنین باید بتوانند در عملیات مداخله کنند یا سیستم

را از طریق دکمه «توقف» یا یک روش مشابه قطع کنند. این بدان معنی است که تعهدات از نوع فنی هستند و بنابراین ارائه‌دهندگان باید اطمینان حاصل کنند که سیستم‌های هوش مصنوعی پرخطر آن‌ها به گونه‌ای طراحی شده است که امکان دخالت انسان را فراهم کند. با این حال، هرگونه قدرت مداخله مستقیماً به ناظران انسانی نسبت داده نمی‌شود. همچنین هیچ تعهد مستقیمی برای کاربران سیستم‌های هوش مصنوعی پرخطر برای اعطای چنین اختیاراتی به ناظران وجود ندارد؛ بنابراین ماده ۱۴ قانون جدید هوش مصنوعی اروپا عمدتاً به جنبه‌های عملکردی و نه اجرایی کنترل انسانی بر سیستم‌های هوش مصنوعی پرخطر توجه دارد. البته به‌زودی به این موضوع بازخواهیم گشت که چگونه برخی از مسئولیت‌های کاربر همچنان به الزامات نظارت در ماده ۱۴ مرتبط است. باتوجه به نوع اطلاعاتی که به‌طور خاص باید در اختیار ناظران انسانی قرار گیرد، مواد ۱۲-۱۳ قانون جدید هوش مصنوعی اروپا از توجه ویژه‌ای برخوردار است، زیرا الزامات ثبت سوابق و شفافیت را از طریق اطلاعات ارائه‌شده به کاربران تعیین می‌کند. به عبارت دیگر، اولاً، ماده ۱۲ قانون جدید هوش مصنوعی اروپا با الزام به ثبت خودکار رویدادها در زمانی که یک سیستم هوش مصنوعی پرخطر در حال کار است، ارائه‌دهندگان را ملزم می‌کند تا اطمینان حاصل کنند که سیستم، سطحی «مناسب» از قابلیت ردیابی عملکرد خود ارائه می‌دهد. هدف موردنظر سیستم به‌طور کلی، حفظ گزارش‌ها در طول عملیات یک سیستم همچون یکی از اقدامات مهم و ممکن برای افزایش شفافیت و به‌ویژه با قابلیت ردیابی فرایندهای سیستم است. با این حال، قانون جدید هوش مصنوعی اروپا در مورد محتوای خاص این مطالب جزئیات چندانی ارائه نمی‌کند که این موضوع میزان ارائه جزئیات را در خود مقررات در مورد اهداف خاص نظارت انسانی محدود می‌کند. براساس ماده ۱۲ باید امکان نظارت بر اینکه آیا سیستم در سطح ملی خطر دارد یا در حال بازسازی است، فراهم شود. همچنین، گزارش‌ها باید نظارت پس‌ازفروش سیستم توسط ارائه‌دهنده را تسهیل کنند، جایی که او موظف است به‌طور فعال و سیستماتیک داده‌های مرتبط ارائه‌شده توسط کاربران یا دیگران را در طول عمر خود جمع‌آوری، مستندسازی و تجزیه و تحلیل کند. این ماده به ارائه‌دهندگانی اشاره می‌کند که هنگام ارائه این قابلیت‌ها، «استانداردهای

شناخته شده» یا «مشخصات مشترک» هوش های مصنوعی پرخطر را رعایت کنند (که ممکن است میان کاربردهای بخش خاص متفاوت باشد). برای ماده ۱۲، هر دو نسخه پیش نویس شورا و پارلمان اروپا جزئیات بیشتری را پیشنهاد می کنند. هرچند این اضافات در وهله اول به اهداف مورد نظر مربوط می شود نه مشخص کردن الزامات ثبت و ورود به سیستم. این اهداف شامل این موارد است که ورود به سیستم باید قابلیت ردیابی را برای شناسایی موقعیت هایی که ممکن است خطری را در مفهوم ماده ۶۵ (۱) ایجاد کند یا منجر به تغییر اساسی شود، امکان پذیر کند. آن ها همچنین به تسهیل نظارت پس از عرضه به بازار همان طور که در ماده ۶۱ اشاره شده است و نظارت بر سیستم های هوش مصنوعی پرخطر در طول عملیات، همان طور که در ماده ۲۹ (۴) اشاره شده است، می پردازند.

باتوجه به ماده ۱۳ قانون جدید هوش مصنوعی اروپا، این مستلزم آن است که سیستم های هوش مصنوعی پرخطر باید طوری طراحی و توسعه داده شوند که عملکرد آن ها به اندازه کافی شفاف باشد تا کاربران بتوانند، خروجی سیستم را تفسیر کرده و به طور مناسب از آن استفاده کنند. درست مانند ماده ۱۴ که بر لزوم شفافیت تأکید می کند. در مورد چگونگی تحقق این هدف، ماده ۱۳ همچنین جزئیاتی را در مورد نوع اطلاعاتی که کاربران سیستم و ناظران انسانی در هنگام نظارت باید به آن دسترسی داشته باشند (و قادر به ارزیابی اش باشند)، ارائه می دهد. این اطلاعات شامل موارد زیر است:

اطلاعات تماس با ارائه دهنده یعنی اطلاعات سیستمی مانند میزان دقت، استحکام و امنیت سایبری، خطرات شناخته شده یا قابل پیش بینی برای سلامت و ایمنی یا حقوق اساسی و اطلاعات مرتبط در مورد مشخصات سیستم برای داده های ورودی؛ بنابراین بر شفافیت از طریق مشخصات فنی که توسط ارائه دهنده «از پیش تعیین شده» هستند، تمرکز می کند. باتوجه به اهمیت نحوه تبدیل اهداف شفافیت ماده ۱۳ قانون جدید هوش مصنوعی اروپا به دامنه کاربر سیستم، ماده ۱۳ (۲) قانون جدید هوش مصنوعی اروپا معتقد است که ارائه دهندگان باید سیستم ها را با دستورالعمل های همراه، توسعه داده و عرضه کنند. اینها باید شامل اطلاعاتی درباره اقدامات نظارتی انسانی و همچنین اقدامات فنی موجود برای کمک به کاربران در تفسیر خروجی های سیستم باشد؛ بنابراین پیوندی که

میان طراحی ارائه‌دهنده و کاربران نهایی این سیستم‌ها در یک سیستم هوش مصنوعی پرخطر ایجاد می‌شود، از طریق ماده ۲۹ (۱) قانون جدید هوش مصنوعی اروپا نیز بیشتر تقویت می‌شود. این ماده کاربران را موظف می‌کند از طریق نظارت بر عملکرد سیستم براساس دستورالعمل چگونگی استفاده، این موارد را دنبال کنند. نسخه پیش‌نویس شورا حتی بیشتر بر پیوند میان دستورالعمل‌های چگونگی استفاده و الزام نظارت بر ماده ۱۴ تأکید می‌کند، زیرا ماده ۲۹ (۴) آن شامل اصلاحیه‌ای پیشنهادی بود که بیان می‌کند، کاربران نیز باید نظارت انسانی را براساس دستورالعمل‌های چگونگی استفاده، اجرا کنند. علاوه بر این، بند ۶ ماده ۲۹ در تمام نسخه‌های پیش‌نویس، تعهد صریح را برای کاربران، به‌ویژه استفاده از اطلاعات ارائه‌شده در ماده ۱۳ برای انجام ارزیابی تأثیر حفاظت از داده‌ها براساس ماده ۳۵ قانون حفاظت از داده‌ها، بیان می‌کند. همان‌طور که برخی تأکید کرده‌اند

(Lazcoz, 2022: 10)، این بدان معناست که کاربران (کنترل‌کنندگان در قانون حفاظت از داده‌ها) موظف هستند از اطلاعات و دستورالعمل‌های ارائه‌دهنده برای فعال کردن افراد (در کنترل) استفاده کنند. در واقع وسعت بالقوه تعهدات کاربر از طریق ماده ۲۹ (۲) قانون جدید هوش مصنوعی اروپا محدود شده است. این ماده در کلیه نظرات کمیسیون، شورا و نسخه‌های پیش‌نویس پارلمان اروپا و نسخه نهایی قانون ۲۰۲۴ به این شکل بیان شده است که الزام به پیروی از دستورالعمل‌ها بدون لطمه به سایر تعهدات کاربر تحت قوانین اتحادیه یا ملی است و شاید مهم‌تر از آن در زمینه نحوه توزیع مسئولیت‌های مرتبط با نظارت انسانی میان ارائه‌دهندگان و کاربران؛ همچنین، این پاراگراف بیان می‌کند که الزام به پیروی از دستورالعمل‌ها بدون لطمه‌ای به اختیار کاربر در سازمان‌دهی منابع و فعالیت‌های خود است (Koulu, 2020: 720)؛ بنابراین قانون جدید هوش مصنوعی اروپا اعتماد زیادی به ارائه‌دهندگان دارد تا کاربران را به اهداف خاص نظارت راهنمایی کنند. به‌طور خلاصه، هدف ماده ۱۴ قانون جدید هوش مصنوعی اروپا ارائه جزئیات زیادی است. در واقع قانون جدید هوش مصنوعی اروپا فضای زیادی را برای ارائه‌دهندگان باقی می‌گذارد تا جزئیات مربوط به آنچه را که به ناظران انسانی ارائه خواهد شد، تعیین کنند.

۳- زمان نظارت انسانی بر هوش مصنوعی توسط ناظران

اما چه زمانی، نظارت انسانی باید انجام شود؟ قانون دریایی ایران در مواد ۶۹ و ۱۱۲، زمان مشخصی برای لزوم نظارت صاحب کشتی یا متصدی حمل بر کارکنان مشخص نمی‌کند و به نظر می‌رسد از زمان استخدام کارکنان تا زمانی که در کشتی حضور دارند، تحت نظارت باید باشند؛ زیرا ماده ۸۲ قانون مذکور به لزوم استخدام افراد مناب اشاره دارد که مؤید لزوم نظارت در زمان استخدام است و ماده ۵۴ قانون دریایی نیز به لزوم تجهیز کارکنان قبل از هر سفر اشاره دارد و ماده ۶۹ و ۱۱۲ نیز به لزوم ادامه نظارت انسانی در زمان انجام وظایف و حضور در کشتی یا به مناسبت فعالیت در کشتی اشاره می‌کند. ماده ۷ قانون مسئولیت مدنی به لزوم نظارت سرپرست یا نگهدارنده صغیر و مجنون اشاره دارد، اما بازه زمانی این نظارت را مشخص نمی‌کند که به نظر می‌رسد در تمام مدتی که به موجب قانون یا قرارداد، صغیر یا مجنون تحت نظارت یا سرپرستی فرد باشد، فرد ناظر یا سرپرست موظف به کنترل و نظارت بر صغیر یا مجنون است. همچنین که ماده ۳۳۳ قانون مدنی تلویحاً به آن اشاره دارد، یعنی تا زمانی که اشیاء تحت مالکیت یک فرد قرار دارد، او موظف به کنترل و نظارت بر آن است؛ اما ماده ۱۲ قانون مسئولیت مدنی معیار خاصی را برای اینکه نظارت باید چه زمانی انجام شود، مشخص می‌کند؛ زیرا بیان می‌دارد که در صورت حادثه‌ای که در زمان انجام کار یا به مناسبت انجام کار، توسط کارگر یا کارمند ایجاد شود، کارفرما مسئول است؛ یعنی نظارت و کنترل کارفرما باید در زمان انجام کار صورت گیرد. از این مواد نمی‌توان استنباط کرد که طول دوره نظارت بر هوش مصنوعی تا چه زمانی است. این نظارت فقط محدود به زمان تولید و عرضه می‌شود یا تا زمانی که به‌روزرسانی‌های هوش مصنوعی به سازنده یا طراح وابسته است. این نظارت باید ادامه پیدا کند؟ اما در قانون جدید اتحادیه اروپا، نظارت چه زمانی بایستی توسط اپراتورها انجام شود؟

در واقع سؤال مهم دیگر برای تعیین محتوای ماهوی نظارت انسانی در ماده ۱۴ قانون جدید هوش مصنوعی اروپا این است که در مورد زمان اجرای نظارت چه چیزی را تجویز می‌کند. در این زمینه، «چه زمانی» در درجه اول به عوامل زمانی صرف اشاره نمی‌کند (اگرچه به‌موقع بودن هرگونه مداخله انسانی موردنیاز، در رابطه با عملیات سیستم هوش

مصنوعی پرخطر مهم است). اما در اینجا، تمرکز بیشتر بر این است که ناظر انسانی در چه مراحل از یک فرایند خودکار باید وظیفه نظارت را انجام دهد. فرایندی که منجر به تصمیم‌گیری یا پشتیبانی از یک سیستم هوش مصنوعی پرخطر می‌شود، می‌تواند از چندین عنصر تشکیل شده باشد. بسته به اینکه چه دیدگاهی اعمال می‌شود، نقطه شروع و پایان این فرایند نظارت می‌تواند، متفاوت باشد. این امر همچنین به این معنی است که «نظارت انسانی» بسته به اینکه در چه مرحله‌ای از فرایند اعمال می‌شود، می‌تواند کارکردهای متفاوتی داشته باشد. برای مثال، آیا نظارت باید روی الگوریتم‌هایی اعمال شود که فرایندهای سیستم را اجرا می‌کنند، یا باید روی خروجی‌های سیستم در قالب توصیه‌ها یا تصمیم‌گیری اعمال شود؟ آیا باید در فواصل منظم و براساس دستورالعمل‌های خاصی نظارت صورت گیرد، یا فقط در پاسخ به برخی از موارد از پیش تعریف‌شده‌ای صورت گیرد که نشان‌دهنده خطر عملکرد اشتباه (داخلی از سیستم یا خارجی از سوی افراد مربوطه) است؟ آیا نظارت باید متوجه دقت کلی و عملکرد صحیح سیستم باشد یا نظارت باید در مقابل دقتی صورت گیرد که سیستم در موارد خاص بایستی انجام دهد، برخورد می‌کند؟ و اگر قرار باشد این تمرکزهای مختلف نظارت با هم ترکیب شوند، در چه پیکربندی باید این کار انجام شود؟

در اینجا نیز، ماده ۱۴ قانون جدید هوش مصنوعی اروپا پاسخ مستقیمی به این موضوع نمی‌دهد که چگونه و با چه معیاری نظارت به مراحل خاصی از یک فرایند سیستم مربوط می‌شود. دو مورد از تعهدات، به صراحت ویژگی «دائمی» دارند، زیرا به عملکرد سیستم «در سراسر صفحه» مربوط می‌شوند؛ بنابراین محدود به زمان نیستند. اینها تعهداتی هستند که در ۱۴(۴) (آ) در مورد ارائه‌دهندگان برای اطمینان از اینکه ناظران قادر به نظارت صحیح بر عملکرد آن هستند و تعهد مندرج در ۱۴(۴) (آ) در مورد ارائه‌دهندگان برای تجهیز سیستم به گونه‌ای است که ناظران می‌توانند، در عملکرد سیستم مداخله کنند یا سیستم را از طریق دکمه «توقف» یا رویه‌ای مشابه قطع کنند. پس نظارت دائمی بوده و در تمام مراحل باید صورت گیرد. علاوه بر این، سه مورد از تعهدات به صراحت به تعامل ناظران با خروجی سیستم مربوط می‌شود. ماده ۱۴ (۴) (ب) ارائه‌دهندگان را ملزم می‌کند که

اطمینان حاصل کنند، سیستم به گونه‌ای طراحی شده است که ناظران انسانی از پیامدهای تکیه بیش از حد به خروجی سیستم آگاه باشند. ماده ۱۴ (۴) (ج)، ایجاب می‌کند که ناظران باید بتوانند، به درستی خروجی سیستم را تفسیر کنند. در نهایت، بند (ب) ماده ۱۴ ایجاب می‌کند که ناظران انسانی باید (از لحاظ فنی) بتوانند، در هر موقعیت خاص تصمیم بگیرند که از سیستم هوش مصنوعی پرخطر استفاده نکنند یا خروجی آن را نادیده بگیرند، یا معکوس کنند. از این نظر، «خروجی‌ها» را نباید فقط به تصمیمات یا توصیه‌های نهایی سیستم تعبیر شوند. ماده ۳ (۱) قانون جدید هوش مصنوعی اروپا، به خروجی‌های سیستم با مثال‌های غیرجامع مانند محتوا، پیش‌بینی‌ها، توصیه‌ها یا تصمیمات اشاره می‌کند؛ بنابراین، آنچه در قامت خروجی واجد شرایط می‌شود، به انتخاب‌های طراحی ارائه‌دهنده بستگی دارد؛ از جمله اینکه سیستم چه داده‌هایی را جمع‌آوری، ثبت، ارزیابی و تولید می‌کند و به راحتی در اختیار ناظران انسانی قرار می‌دهد. با یادآوری بحث در بخش قبل، در مورد تعهدات کاربران برای پیروی از دستورالعمل‌های ارائه‌دهندگان، ارائه‌دهنده مجاز است جزئیات الزام‌آوری را به دستورالعمل‌ها اضافه کند؛ از جمله اینکه در یک انسان باید در چه مراحل از فرایند سیستم یا فرایندهای ترکیبی، نظارت خود را اعمال کند.

برای اطمینان از اینکه نظارت انسانی واقعاً توسط کاربر انجام می‌شود، ارائه‌دهندگان سیستم‌های پرخطر می‌توانند، فرایندهای سیستم را به صورت ترکیبی یا نیمه خودکار طراحی کنند. آن‌ها می‌توانند، این کار را با الزام به دخالت انسانی در مراحل خاصی از فرایند سیستم انجام دهند. برای مثال در برنامه‌های «ردکننده»^۱ سیستم تصمیم می‌گیرد

۱. برنامه‌های «ردکننده» یا «gatekeeping programs» به نرم‌افزارها یا سیستم‌هایی اشاره دارند که تصمیم می‌گیرند آیا یک وظیفه یا فرایند باید توسط سیستم خودکار انجام شود یا به متخصص انسانی ارجاع داده شود. این سیستم‌ها نقش کلیدی در تعیین مسیر کارها و فرایندها دارند و معمولاً در زمینه‌هایی مانند مدیریت محتوای آنلاین، فرایندهای اداری و سایر سیستم‌های تصمیم‌گیری مورد استفاده قرار می‌گیرند. هدف این برنامه‌ها اطمینان از این است که وظایف به بهترین نحو ممکن انجام شود، خواه توسط هوش مصنوعی یا نیروی انسانی. به عنوان مثال: -در برنامه‌های مدیریتی محتوا، سیستم ممکن است تصمیم بگیرد که یک مطلب به صورت خودکار تأیید و منتشر شود یا برای بازبینی انسانی ارسال شود.

که آیا یک وظیفه معین به‌بهرترین نحو توسط سیستم یا یک متخصص انسانی انجام شده است یا نه. ماده ۱۴ قانون جدید هوش مصنوعی اروپا به چنین مواردی از عملکرد نظارتی اشاره نکرده است، به غیر از مواردی که از سیستم هوش مصنوعی برای شناسایی بیومتریک و طبقه‌بندی اشخاص حقیقی استفاده می‌شود. با این حال، برای تمام سیستم‌های پرخطر دیگر، این ماده فقط قابلیت‌های فنی را برای ناظران برای مداخله فعال بر رویه و نادیده گرفتن خروجی سیستم تجویز می‌کند. لازم به ذکر است که این شرایط به‌هیچ‌وجه مانع از آن نمی‌شود که ارائه‌دهندگان به سیستم‌های پرخطر خود دستور دهند تا فرایندی را در موارد یا مراحل از پیش تعیین شده به سمت نظارت انسانی هدایت کنند. بسیاری از سیستم‌های هوش مصنوعی شامل تصمیم‌های طراحی ثابت در مورد نیاز به ورودی انسانی هستند، مانند سیستم‌های توصیه‌گر یا جایی که سیستم‌ها برای قطع فرایندها براساس مشکلات در تفسیر داده‌های ورودی برنامه‌ریزی شده‌اند.^۱ این نوع

-در خدمات مشتری، سیستم می‌تواند تصمیم بگیرد که یک درخواست مشتری توسط ربات پاسخ داده شود یا به متخصص انسانی منتقل شود.

۱ سیستم‌های توصیه‌گر و سیستم‌هایی که برای قطع فرایندها بر اساس مشکلات در تفسیر داده‌های ورودی برنامه‌ریزی شده‌اند، در حوزه‌های مختلف کاربرد دارند. در زیر به چند مثال از هر دو نوع سیستم اشاره می‌شود: سیستم‌های توصیه‌گر (Recommender Systems)

۱. سیستم‌های توصیه‌گر فیلم و موسیقی:

۵ تئفلیکس: تئفلیکس از سیستم‌های توصیه‌گر استفاده می‌کند تا بر اساس تاریخچه تماشای کاربر، پیشنهاد‌های مناسبی برای فیلم‌ها و سریال‌ها ارائه دهد.

۵ اسپاتیفای: اسپاتیفای بر اساس سلیقه موسیقی کاربر و تاریخچه گوش دادن به موسیقی، لیست‌های پخش پیشنهاد می‌دهد.

۲. فروشگاه‌های آنلاین:

۵ آمازون: آمازون از سیستم‌های توصیه‌گر استفاده می‌کند تا بر اساس خریدهای قبلی و مرور محصولات توسط کاربر، کالاهای مشابه و مرتبط را پیشنهاد دهد.

سیستم‌هایی برای قطع فرایندها بر اساس مشکلات در تفسیر داده‌های ورودی

۱. سیستم‌های پردازش تراکنش‌های بانکی:

۵ تشخیص تقلب: بانک‌ها از سیستم‌های هوش مصنوعی برای مانیتورینگ تراکنش‌های مالی استفاده می‌کنند. اگر سیستم به یک تراکنش مشکوک شود (مثلاً مبلغ غیرمعمول یا محل جغرافیایی غیرمنتظره)، تراکنش را متوقف کرده و برای بررسی به یک متخصص انسانی ارجاع می‌دهد.

ترتیبات فنی برای حصول اطمینان از اعمال نظارت انسانی، همچنان قانونی و در بسیاری موارد ارجح خواهد بود. بااین حال، عبارت ماده ۱۴ تاحدزیادی مشخص نمی کند که چه زمانی یک سیستم باید یک فرایند را به نظارت انسانی هدایت کند و به ناظر انسانی ارجاع دهد. این شرایط، خطر کاهش تأثیر مقررات را بر گستره ای که نظارت انسانی واقعاً انجام می شود، همراه دارد. در مجموع، ماده ۱۴ قانون جدید هوش مصنوعی اروپا و همچنین الزامات مندرج در ماده ۱۳ قانون جدید هوش مصنوعی اروپا در مورد محتوای خاصی که ارائه دهنده باید در دستورالعمل های خود به کاربر درج کند، راهنمایی نسبتاً مبهمی در مورد اینکه «چه زمانی» نظارت انسانی بر سیستم های هوش مصنوعی پرخطر باید انجام شود، ارائه می دهد. این امر، به ویژه در مورد نوع انگیزه هایی که نیاز به نظارت دارند، صادق است. کاربران سیستم های هوش مصنوعی پرخطر نیز ممکن است، مشمول مقرراتی باشند که تعهدات برای استفاده از قابلیت های نظارت داخلی در مراحل خاصی از یک فرایند سیستم را اعمال می کنند. البته ممکن است، کاربران بخواهند بر جنبه های حساس آن وظایفی که به سیستم سپرده شده است، کنترل داشته باشند. اینکه ماده ۱۴ قانون جدید هوش مصنوعی اروپا، روشن می کند که کاربران در هر مورد باید ظرفیت فنی برای قطع فرایند را داشته باشند و اینکه آن ها باید بتوانند، خروجی سیستم را نادیده بگیرند؛ بنابراین این امر یک جزء مهم در تضمین بازمیانی انسانی از ویژگی های اساسی است. درعین حال، تمرکز بر توانایی کاربر برای مداخله فعال به طور کامل این واقعیت را نشان نمی دهد که ناظر انسانی خاص ممکن است در شناسایی مراحل از فرایند که از ورودی های انسانی سود می برند، بدون کمک خود سیستم، با مشکل مواجه شود؛ بنابراین، اگرچه قانون جدید هوش مصنوعی اروپا اطمینان حاصل می کند که کاربر می تواند هر زمان که بخواهد ورودی انسانی را در فرایند وارد کند، انتخاب های طراحی

۲. سیستم های مدیریت تولید و کیفیت:

0 کارخانه های هوشمند: در کارخانه های هوشمند، سیستم های نظارت بر کیفیت تولید، داده های ورودی از سنسورها را تحلیل می کنند. اگر مشکلی در خط تولید شناسایی شود (مثلاً نقص در یک محصول)، سیستم می تواند خط تولید را به طور خودکار متوقف کرده و پرسنل فنی را برای بررسی و رفع مشکل مطلع کند.

ارائه‌دهنده، تأثیر زیادی بر تعیین زمان مداخله کاربر دارد و در مورد چگونگی اطمینان از اینکه قابلیت‌های نظارتی ماده ۱۴ قانون جدید هوش مصنوعی اروپا اعمال شده است، تأثیر خواهد داشت.

۴- شناسایی ناظران انسانی

در حقوق ایران در ماده ۶۹ قانون دریایی، صاحب کشتی، در قامت مسئول نظارت معین شده است و در ماده ۱۱۲ متصدی حمل را مسئول نظارت می‌داند. در ماده ۱۲ قانون مسئولیت مدنی، کارفرما و ماده ۷ سرپرست مجنون یا صغیر و در ماده ۳۳۳ قانون مدنی مالک اشیا مسئول نظارت معرفی شده‌اند؛ بنابراین جمع این مواد ملاک واحدی برای تعیین اینکه چه کسی باید نظارت کند، به ما نمی‌دهد. در صورت اشیا تلقی کردن هوش مصنوعی، شاید بتوان از ماده ۳۳۳ قانون مدنی وحدت ملاک گرفت و مالک هوش مصنوعی را مسئول نظارت تلقی کرد. اما مالک هوش مصنوعی چه کسی است؟ کسی که از آن استفاده می‌کند یا کسی که آن را تولید و طراحی کرده است یا کسی که مسئول به‌روزرسانی و ارتقای آن است؟ بنابراین از این ماده نمی‌توان برای شناسایی مالک هوش مصنوعی استفاده کرد و اگر برای هوش مصنوعی، شخصیت حقوقی قائل شویم، به‌خصوص در هوش مصنوعی کمک دیجیتال، همچون کارمند می‌تواند محسوب شده و مشمول ماده ۱۲ قانون مسئولیت مدنی شده و فردی که آن را به کار گرفته است را مسئول نظارت بر هوش مصنوعی دانست. اما در قانون ۲۰۲۴ اروپا چه کسی در قامت مسئول نظارت معرفی شده است؟

«چه کسی باید نظارت انسانی را انجام دهد؟» سؤالی است که اهمیت عملی زیادی برای کیفیت نظارتی دارد که توسط انسان باید انجام شود و در نتیجه برای کارایی آن همچون یک اقدام حفاظتی نیز از اهمیت بالایی برخوردار است. این امر به این دلیل است که شایستگی (دانش) و همچنین اختیارات (قدرت) ناظران انسانی در دو مورد تأثیر می‌گذارد - ۱- بر انتظارات معقول در مورد آنچه که آن‌ها قادر به تشخیص و واکنش به آن هستند و ۲- در مورد نوع اقداماتی که می‌توانند، انجام دهند. ماده ۱۴ قانون جدید هوش مصنوعی اروپا بر جنبه‌های ارتباطی انتقال دانش میان سیستم‌ها و انسان‌ها تأکید می‌کند.

بااین حال، به طور غیرمستقیم و کلی، قانون جدید هوش مصنوعی اروپا نشان می دهد که سطح مشخصی از دانش، از آن انسان هایی انتظار می رود که توسط کاربر وظیفه انجام نظارت را دارند. برای مثال، بند ۴۸ بیان می کند که اقدامات نظارتی انسانی، در صورت لزوم، باید تضمین کند که سیستم مجهزه فاکتورهایی است که به اپراتور انسانی پاسخ می دهد و اینکه افراد حقیقی کسانی که نظارت انسانی به آن ها محول شده است، شایستگی، آموزش و اختیار لازم برای انجام آن نقش را دارند. از آنجایی که این تعهدات تحت اختیار کاربر مسئول اشاره می کند، دشوار است که مشخص کنیم، ارائه دهندگان چگونه می توانند چنین تضمین هایی را فراهم کنند. این امر که کاربران چگونه کار خود را سازماندهی می کنند، معمولاً خارج کنترل آن ها است. به عبارت دیگر، از ماده ۹ (۴) قانون جدید هوش مصنوعی اروپا مشخص است که هیچ «ضمانتی» نباید توسط ارائه دهندگان ارائه شود. این ماده صرفاً ارائه دهندگان را موظف می کند که به دانش فنی، تجربه و آموزش مورد انتظار کاربر و محیطی که قرار است، سیستم در آن استفاده شود، توجه لازم را داشته باشند.

به طور خلاصه، از آنجایی که ماده ۱۴ قانون جدید هوش مصنوعی اروپا خطاب به ارائه دهندگان سیستم است، به خودی خود به این موضوع اشاره نمی کند که به طور خاص چه کسی در پایان کار باید نظارت را انجام دهد. ارائه دهنده باید دانش کاربران مورد نظر را در هنگام طراحی قابلیت های نظارت فنی (و همچنین هنگام تنظیم دستورالعمل های کاربر همراه) در نظر داشته باشد. فقدان راهنمایی خاص در مورد تعهدات ارائه دهنده در این زمینه، ارزیابی دقیق بودن این ماده را دشوار می کند. با وجود این، ملاحظات و انتظارات ارائه دهنده از دانش، تجربه و آموزش کاربر برای کاربر الزام آور نخواهد بود.

به نظر جزئیات نظارتی نسبتاً کم قانون جدید هوش مصنوعی اروپا در مورد مسائل مربوط به صلاحیت در نظارت انسانی، خطر و نقص مقررات است. بااین حال، در اینجا باید اضافه شود که نسخه پیش نویس شورا شامل یک تعهد صریح جدید برای کاربران است که «نظارت انسانی را به اشخاص حقیقی که صلاحیت، آموزش و اختیار لازم را دارند واگذار کنند». به طور مشابه، در نسخه پیش نویس پارلمان اروپا، ماده ۲۹ (۱) (آ) و (۲) جدید، کاربران را

موظف می‌کند، اطمینان حاصل کنند که «افراد حقیقی که برای اطمینان از نظارت انسانی بر سیستم‌های هوش مصنوعی پرخطر تعیین شده‌اند، صلاحیت دارند و واجد شرایط مناسب هستند و آموزش‌دیده» باشند، و همچنین ناظران «منابع لازم را برای تضمین نظارت مؤثر بر سیستم هوش مصنوعی مطابق با ماده ۱۴» در اختیار داشته باشند. برخی از انواع این تعهدات در شکل نهایی قانون جدید هوش مصنوعی اروپا منعکس شد و پیوند میان تعهدات ارائه‌دهنده در ماده ۱۴ و تعهدات کاربر در ماده ۲۹ قانون جدید هوش مصنوعی اروپا را بیشتر تقویت کرد؛ بنابراین چنین تعهداتی الزامات اساسی بیشتری را برای کاربران ایجاد می‌کند. درنهایت، کیفیت نظارت به این سؤال اساسی مربوط می‌شود که آیا ورودی‌های انسانی انجام شده است یا خیر. در رابطه با قانون حفاظت از داده‌ها، هیئت حفاظت از داده‌های اروپا و دفتر کمیساریای اطلاعات بریتانیا^۱، تأکید کرده‌اند که آموزش کارکنانی که باید نظارت انسانی بر سیستم‌های خودکار را انجام دهند، برای اطمینان از اینکه سیستم تا حدی خودکار در نظر گرفته می‌شود، انجام خواهد شد و مشروط به الزامات خاص تعیین‌شده برای فرایندهای تصمیم‌گیری صرفاً خودکار، که الزامات تعیینی در ماده ۲۲ قانون حفاظت از داده‌ها لازم باشد. مهم این است که فراموش نکنیم که این سؤال که نظارت باید «توسط چه کسی» انجام شود، فقط یکی از تقسیم‌بندی‌های رسمی مسئولیت‌ها نیست. هنگام ارزیابی کیفیت عملکرد نظارت، تمرکز باید تا حدی از اقدامات و مسئولیت‌های کاربر به کارمندان یا کارمندان دولتی که قرار است با سیستم پرخطر تعامل داشته باشند، منتقل شود. اینکه این تعامل امکان ارزیابی وزمینه را فراهم می‌کند و میان نبود تقارن انسان و ماشین پل می‌زند، اساسی است. درحالی‌که نظارت انسانی «معنادار»، موضوعی بسیار پیچیده‌تر است که از حوصله این مقاله خارج

1 United Kingdom Information Commissioner's office, 'How do we ensure individual rights in our AI systems?', <<https://ico.org.uk/for-organisations/guide-to-data-protection/key-dp-themes/guidance-on-ai-and-data-protection/how-do-we-ensure-individual-rights-in-our-ai-systems/>>, accessed 9 September 2022. See, also, European Data Protection Board and European Data Protection Supervisor, 'Joint Opinion 5/2021 on the proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)' 5/2021 18 June 2021, p. 6, at point 7.

است و همچنین به احتمال زیاد پیچیده تر از آن است که بتوان آن را از طریق چند مفاد که مجموعه‌ای از سیستم‌های هوش مصنوعی را هدف قرار می‌دهند، حل کرد. در حقوق ایران، مسئولیت‌های نظارتی به طور پراکنده در مواد مختلف قانون تعریف شده‌اند، اما یک قانون جامع و مدون برای تعیین ناظر انسانی بر هوش مصنوعی وجود ندارد. در اینجا می‌توان از ماده ۳۳۳ قانون مدنی برای تشخیص مالک استفاده کرد و ممکن است این مالک در قامت مسئول نظارت معرفی شود، اما همان‌طور که اشاره شد، مسائل مربوط به «مالکیت» هوش مصنوعی پیچیده است.

از آنجایی که در قانون مسئولیت مدنی ایران، کارفرما یا صاحب‌کار همچون ناظر بر کارهای کارکنان یا اشیاء شناخته شده است، می‌توان به طور فرضی مسئولیت نظارت بر هوش مصنوعی را به مالک یا کاربر سیستم واگذار کرد. در صورتی که هوش مصنوعی شخصیت حقوقی داشته باشد، مانند هوش مصنوعی‌های دیجیتال یا کارمندان دیجیتال، مسئولیت نظارت ممکن است به کارفرما یا مسئول سازمان منتقل شود. پس پیشنهاد می‌شود:

۱- واضح‌تر کردن مفاهیم و تعاریف قانونی: در قوانین داخلی مانند ایران، باید تعریف دقیق‌تری از «مالک» هوش مصنوعی و نقش‌های نظارتی در زمینه‌های مختلف (کشاورزی، حمل‌ونقل، پزشکی و...) ارائه شود.

۲- ایجاد چارچوب آموزشی و نظارتی: برای ناظران انسانی، باید تدابیر مشخصی برای آموزش و نظارت وجود داشته باشد تا از کیفیت و اثربخشی نظارت اطمینان حاصل شود. برای مثال، ناظران باید درک کاملی از خطرات و مزایای سیستم‌های هوش مصنوعی و توانایی تعامل و مداخله در شرایط ضروری را داشته باشند.

۳- اهمیت نظارت چندسطحی: برای افزایش کارایی نظارت ممکن است، بهتر باشد که مسئولیت نظارت فقط به یک فرد محدود نشود و افراد مختلف در سطوح مختلف نظارتی، از توسعه‌دهندگان گرفته تا کارمندان اجرایی و ناظران مستقل، در این فرایند دخیل باشند. این نظارت باید به طور مستمر و با توجه به پیشرفت‌های فناوری هوش مصنوعی به‌روزرسانی شود. در غیراین صورت، ممکن است، قوانین نتوانند به طور مؤثر به پیشرفت‌ها و چالش‌های سریع این حوزه پاسخ دهند.

۵- مبنای مسئولیت ناشی از نظارت بر هوش مصنوعی

در رابطه با مسئولیت نظارت بر هوش مصنوعی، سؤال اصلی این است که مسئولیت در صورت بروز خطا یا آسیب از کجا آغاز می‌شود و تا چه حد گسترده است. مسئولیت می‌تواند چندلایه داشته باشد:

۱. **مسئولیت توسعه‌دهنده یا سازنده سیستم هوش مصنوعی:** توسعه‌دهندگان باید اطمینان حاصل کنند که سیستم‌های هوش مصنوعی طراحی شده به‌طور ایمن و مطابق با استانداردهای اخلاقی عمل کنند. در صورت بروز مشکلات ناشی از نقص در طراحی یا توسعه، ممکن است، مسئول شناخته شوند.

۲. **مسئولیت اپراتورها و نظارت‌کنندگان انسانی:** اگر سیستم‌های هوش مصنوعی به‌طور نادرست عمل کنند یا باعث آسیب شوند، مسئولیت نظارت به‌طور ویژه بر دوش اپراتورهای انسانی قرار دارد. در صورتی که نظارت کافی نباشد یا نظارت انسانی به‌درستی اعمال نشود، ممکن است، مسئولیت ناشی از آسیب‌ها به‌عهده اپراتور یا سازمان‌های نظارتی باشد.

۳. **مسئولیت قانون‌گذاران و نهادهای نظارتی:** در مواردی که قوانین و مقررات به‌درستی وضع نشده یا نهادهای نظارتی به‌طور مؤثر نظارت نمی‌کنند، مسئولیت ممکن است، متوجه دولت‌ها یا نهادهای نظارتی باشد.

درنهایت، **قلمرو مسئولیت** در این زمینه بستگی به میزان تأثیر و خطرات ناشی از سیستم‌های هوش مصنوعی خواهد داشت. برای مثال، در سیستم‌های پرخطر و حساس مانند سیستم‌های بهداشتی، حمل‌ونقل یا نظامی، مسئولیت نظارتی می‌تواند، بسیار گسترده‌تر و دقیق‌تر باشد.

همچنین، در زمینه اعتماد بیش‌ازحد به ارائه‌دهندگان سیستم‌های هوش مصنوعی، یکی از چالش‌های اصلی این است که احتمال دارد، نظارت کافی در برابر فناوری‌های پیچیده و خودکار به وجود نیاید. این اعتماد، می‌تواند مشکلاتی را در مواجهه با بحران‌های احتمالی ایجاد کند و اهمیت نظارت انسانی را بیشتر نمایان کند.

پس در سیستم‌های هوش مصنوعی، مانند خودروهای تمام خودران که دیگر شخصی تحت عنوان راننده وجود ندارد، مسئولیت در صورت وقوع حوادث متوجه دست‌اندرکاران دیگر در خودروی تمام خودران خواهد بود از جمله سازنده خودروی تمام خودران در قامت مسئول اصلی، باتوجه به قوانین حمایت از حقوق مصرف‌کنندگان کالا و خدمات و قانون حمایت از مصرف‌کنندگان خودرو در طول تضمین، مبنای مسئولیت سازنده مسئولیت محض بوده و پس از پایان دوره تضمین باید براساس قواعد عام مسئولیت مدنی به او رجوع شود اما در حقوق اتحادیه اروپا مسئولیت سازنده، یک مسئولیت محض است. از سوی دیگر، قانون‌گذار ایران مانند قانون اتحادیه اروپا مسئولیت سازنده را محدود به زمان عرضه به بازار دانسته و مسئولیتی برای پس از به‌گرددش درآوردن خودرو از لحاظ نظارت بر خودرو و انجام به‌روزرسانی‌های لازم، تعیین نکرده است؛ بنابراین مسئولیت او به دلیل فقدان نص قانونی در دوره پس از عرضه، دیگر مبتنی بر مسئولیت محض نبوده و باتوجه به قواعد حاکم بر مسئولیت مدنی با او رفتار شده و مسئولیت مبتنی بر تقصیر برای او وجود خواهد داشت. در رابطه با مسئولیت ناظران فنی نیز، آنها فقط در صورتی مسئول خواهند بود که در صورت هشدار خودروی تمام خودران مداخله نکرده و سبب وقوع حادثه‌ای شوند. مسئولیت اپراتورهای پشتیبان نیز در حقوق ایران مطابق با اصل، یعنی مسئولیت مبتنی بر تقصیر خواهد بود. گرچه در حقوق اروپا مسئولیت محض آنان پذیرفته شده است. اما مبنای مسئولیت ارائه‌دهندگان خدمات، مانند خدمات فناوری مطابق قانون حمایت از مصرف‌کنندگان کالا و خدمات مصوب ۱۳۸۸ پیش‌میان‌ی شده است که براساس قانون و عرف باید سنجیده شود؛ بنابراین با استناد به ماده ۱ قانون مسئولیت مدنی و قواعد قانون مدنی در باب تصویب مبنای مسئولیت ارائه‌دهنده خدمات فناوری مطابق اصل مسئولیت مبتنی بر تقصیر سنجیده شده و مسئولیت محض جنبه استثنا داشته و قابل تفسیر به صورت موسع و تسری به سایر موارد مشابه نیست. گرچه برای حمایت هرچه بیشتر از زیان‌دیده و باتوجه به ماهیت خاص ارائه خدمات فناوری لازم است، مسئولیت محض ارائه‌دهنده خدمت پذیرفته شود، اما با وجود این مطابق قانون حمایت از مصرف‌کننده مبتنی بر تقصیر است.

نتیجه‌گیری

در این مقاله، بر قوانین ایران و قانون جدید هوش مصنوعی اروپا ۲۰۲۴ تمرکز شده است. تجزیه و تحلیل و بررسی‌ها نشان داد که نه ماده ۱۴ قانون جدید هوش مصنوعی اروپا و نه قوانین ایران، جزئیات زیادی در مورد آنچه که ناظر انسانی در هنگام اجرای نظارت باید در نظر بگیرد یا توجه خود را به آن معطوف کند، ارائه نمی‌دهد. همچنین نشان داد که فضای زیادی برای ارائه‌دهندگان باقی خواهد ماند تا جزئیات مربوط به چه داده‌هایی را به ناظران انسانی ارائه کنند، تعیین کند. نتیجه‌گیری‌های مشابهی در مورد «زمان» اعمال نظارت انجام شد. ارائه‌دهندگان باید اطمینان حاصل کنند که کاربران در هر زمان ممکن است، فرایندهای سیستم را قطع یا لغو کنند. با این حال، فقدان راهنمایی در مورد اینکه چه نوع انگیزه‌هایی باعث تعهد به انجام چنین نظارتی می‌شود، به این معنی است که انتخاب‌های طراحی سیستم ارائه‌دهنده، همچنان تأثیر زیادی در زمانی خواهد داشت که نظارت انسانی انجام می‌شود؛ بنابراین، این شرایط نفوذ تعهدات نظارتی قانون جدید هوش مصنوعی اروپا را در استفاده روزمره از این سیستم‌ها محدود می‌کند و بر مقیاسی تأثیر می‌گذارد که در آن ناظران انسانی قادر به کشف یا کاهش خطرات خطاها در موارد مختلف خواهند بود. علاوه بر این، قانون جدید هوش مصنوعی اروپا به این موضوع اشاره نمی‌کند که «توسط چه کسی» باید نظارت را در پایان کار انجام دهد، اگرچه تأکید بر اینکه نسخه‌های پیش‌نویس شورا و پارلمان اروپا بر تعهدات کاربران برای اطمینان از آموزش و صلاحیت‌ها و همچنین منابع ناظران انسانی گذاشته‌اند، تشخیص اهمیت و اولویت این سؤال را برجسته می‌کند. با این حال، در مجموع، اگرچه ارائه‌دهندگان باید طراحی فنی و همچنین دستورالعمل‌ها را برای کاربران سیستم موردنظر اتخاذ کنند، اما کاربران سیستم موردنظر از اختیارات نامحدودی در نحوه انجام نظارت برخوردارند.

منابع و مأخذ

- ۱- خوئی، سید محمد (۱۳۹۶) **هوش مصنوعی و قانونگذاری ۱**، مرکز پژوهش‌های مجلس شورای اسلامی.
- ۲- رجبی، عبدالله (۱۳۹۸) **ضمان در هوش مصنوعی**، مطالعات حقوق تطبیقی، دوره ۱۰، شماره ۲، ۴۴۹-۴۶۶.
- ۳- قاسمی، محمدرضا (۱۴۰۰) **هوش مصنوعی و حکمرانی آینده**، نشریه حوزه، دوره ۳۸، شماره ۱۲، ۱۶۶-۱۷۹.
- 1- Binns, 'Human Judgment in Algorithmic Loops: Individual Justice and Automated Decision-Making' (2022) 16 Regulation & Governance 197
- 2- Enqvist Lena (2023) 'Human oversight' in the EU artificial intelligence act: what, when and by whom?, Law, Innovation and Technology, 15:2, 508-535, DOI: 10.1080/17579961.2023.2245683
- 3- Enarsson, L. Enqvist, and M. Naartijarvi, 'Approaching the Human in the Loop - Legal Perspectives on Hybrid Human/Algorithmic DecisionMaking in Three Contexts' (2022) 31 Information & Communications Technology Law 123.
- 4- Green, 'The Flaws of Policies Requiring Human Oversight of Government Algorithms' (2022) 45 Computer Law & Security Review 1;
- 5- Lazcoz and P. De Hert, 'Humans in the GDPR and AIA Governance of Automated and Algorithmic Systems. Essential Pre-Requisites against Abdicating Responsibilities' (2022) 8 Brussels Privacy Hub Working Paper, p. 10.
- 6- Harasimiuk and T. Braun, **Regulating Artificial Intelligence: Binary Ethics and the Law** (Routledge Research in the Law of Emerging Technologies, Routledge, 2021), 49
- 7- R. Koulu, 'Proceduralizing Control and Discretion: Human Oversight in Artificial Intelligence Policy' (2020) 27 Maastricht Journal of European and Comparative Law 720.
- 8- Methnani and others, 'Let Me Take Over: Variable Autonomy for Meaningful Human Control' (2021) 4 Frontiers in Artificial Intelligence 737072-737072
- 9- Onitiu, 'The Limits of Explainability & Human Oversight in the EU Commission's Proposal for the Regulation on AI- a Critical Approach Focusing on Medical Diagnostic Systems' (2022) 32 Information & Communications Technology Law 170, 181
- 10- Smuha, 'From a "race to AI" to a "race to AI regulation": regulatory competition for artificial intelligence' (2021) 13 Law, Innovation and Technology 57.