



ORIGINAL RESEARCH PAPER

Legal Liability of Artificial Intelligence Algorithms in Decision-Making and the Legal Status of Users

Zahra Salehi ¹, Aida Mokhtare^{*2}

Received:

15 Dec 2025

Revised:

21 Feb 2026

Accepted:

19 Apr 2026

Available Online:

23 Sep 2026

Keywords:

Artificial Intelligence,
Civil Liability, User
Rights, Algorithmic
Decisions, Legal Rules.

Abstract

The rapid expansion of artificial intelligence (AI) systems in sensitive decision-making processes has created fundamental challenges in the fields of civil liability and the legal status of users. This study aims to examine the foundations and challenges of attributing liability arising from algorithmic decisions and to assess the adequacy of traditional legal rules in addressing such challenges. The research employs a descriptive-analytical methodology based on library research and a comparative examination of the regulations of leading legal systems, particularly those of the European Union. The findings indicate that traditional legal concepts such as fault, causation, and foreseeability are no longer fully capable of accurately explaining liability in the context of learning technologies. Without reconsidering these concepts, the possibility of providing effective compensation to injured parties may be seriously limited. The results further demonstrate that users, as stakeholders in AI-based decision-making processes, require recognition of fundamental rights, including the right to information, the right to contest decisions, and the right to human review, in order to establish a legal balance between users and service providers. The comparative analysis shows that the development of binding and structured legal rules, together with the adoption of a multilayered approach based on institutional responsibility, protective mechanisms, and international coordination, is an unavoidable necessity. The overall conclusion of the study is that only through such a framework can justice, transparency, and effective compensation be ensured in the digital environment.

1 M.A. Student in Financial and Economic Law, Apadana Institute of Higher Education, Shiraz, Iran.

2* Assistant Professor and Faculty Member, Apadana Institute of Higher Education, Shiraz, Iran. (Corresponding Author) Email: dr.mokhtare@apadana.ac.ir

Please Cite This Article As: Salehi, Z & Mokhtare, A (2026). "Legal Liability of Artificial Intelligence Algorithms in Decision-Making and the Legal Status of Users". *Interdisciplinary Legal Research*, 7(3): 21-36.



This is an open access article distributed under the terms of the Creative Commons Attribution (CC BY 4.0)

مقاله پژوهشی

(صفحات ۳۶-۲۱)

مسئولیت حقوقی الگوریتم هوش مصنوعی در تصمیم‌گیری‌ها و جایگاه حقوقی کاربر

زهرا صالحی^۱، آیدا مخترع^{۲*}

۱. دانشجوی کارشناسی ارشد حقوق مالی اقتصادی موسسه آپادانا، شیراز، ایران.

۲. استادیار و عضو هیأت علمی موسسه آموزش عالی آپادانا، شیراز، ایران. (نویسنده مسؤول)

دریافت: ۱۴۰۴/۰۹/۲۴ ویرایش: ۱۴۰۴/۱۲/۰۲ پذیرش: ۱۴۰۵/۰۱/۳۰ انتشار: ۱۴۰۵/۰۷/۰۱

چکیده

گسترش روزافزون سامانه‌های هوش مصنوعی در فرایندهای تصمیم‌گیری حساس، چالش‌های بنیادینی را در حوزه مسئولیت مدنی و جایگاه حقوقی کاربران ایجاد کرده است. این پژوهش با هدف واکاوی مبانی و چالش‌های انتساب مسئولیت ناشی از تصمیمات الگوریتمی و ارزیابی کفایت قواعد سنتی حقوقی انجام شده است. روش تحقیق توصیفی - تحلیلی و مبتنی بر منابع کتابخانه‌ای و بررسی تطبیقی مقررات نظام‌های حقوقی پیشرو، به‌ویژه اتحادیه اروپا، است. یافته‌ها نشان می‌دهد که مفاهیم کلاسیکی همچون تقصیر، رابطه سببیت و پیش‌بینی‌پذیری دیگر توان تبیین دقیق مسئولیت در بستر فناوری‌های یادگیرنده را ندارند و بدون بازنگری در این مفاهیم، امکان جبران خسارت زیان‌دیدگان به‌طور جدی محدود می‌شود. نتایج همچنین آشکار ساخت که کاربران، به‌عنوان طرفین ذی‌نفع در فرایند تصمیم‌گیری هوش مصنوعی، نیازمند شناسایی حقوق بنیادینی مانند حق آگاهی، حق اعتراض و حق بازبینی انسانی هستند تا تعادل حقوقی میان آنان و ارائه‌دهندگان خدمات برقرار شود. مطالعه تطبیقی نشان داد که تدوین قواعد الزام‌آور و ساختارمند بهره‌گیری از رویکرد چندلایه مبتنی بر مسئولیت نهادی، سازوکارهای حمایتی و هماهنگی بین‌المللی ضرورتی اجتناب‌ناپذیر است. نتیجه کلی پژوهش آن است که تنها از رهگذر چنین چارچوبی می‌توان عدالت، شفافیت و کارآمدی جبرانی را در فضای دیجیتال تضمین کرد.

کلمات کلیدی: هوش مصنوعی، مسئولیت مدنی، حقوق کاربران، تصمیمات الگوریتمی، قواعد حقوقی.

مقدمه

ورود سامانه‌های هوش مصنوعی به حوزه‌های تصمیم‌گیری حساس، ساختار سنتی مسئولیت مدنی را در معرض دگرگونی قرار داده است. این سامانه‌ها با اتکا بر مدل‌های یادگیرنده و فرآیندهای خودکار، تصمیماتی تولید می‌کنند که منطق درونی آن‌ها اغلب قابل توضیح نیست. همین ویژگی باعث شده است که معیارهای کلاسیکی مانند تقصیر و پیش‌بینی‌پذیری دیگر توان تبیین مسئولیت را نداشته باشند. طراحان نیز به دلیل ماهیت پویا و خودآموز الگوریتم‌ها، کنترل کامل بر خروجی‌ها ندارند. این وضعیت موجب می‌شود نسبت میان داده‌های ورودی و نتیجه نهایی، روشن و قابل بازسازی نباشد. چنین ابهامی مانع از اعمال قواعد سنتی رابطه سببیت می‌گردد. پیامد آن، دشواری تعیین فاعل حقوقی مسؤول در فرآیند تصمیم‌گیری است. این شرایط ضرورت بازاندیشی در مفاهیم بنیادین مسئولیت را برجسته می‌کند. بنابراین، مسأله انتساب مسئولیت در تصمیمات الگوریتمی یکی از چالش‌های اصلی حقوق فناوری شده است.

در کنار چالش‌های مفهومی، جایگاه حقوقی کاربران در مواجهه با سامانه‌های هوش مصنوعی نیز با ابهام جدی روبه‌رو است. بسیاری از کاربران بدون آگاهی کافی از نحوه عملکرد سامانه‌ها، ناخواسته تحت تأثیر تصمیمات خودکار قرار می‌گیرند. این ناآگاهی باعث تضعیف رضایت آگاهانه و برهم خوردن توازن قراردادی می‌شود. برخی شرکت‌های فناورانه نیز با درج شروط محدودکننده، می‌کوشند مسئولیت‌های قانونی خود را کاهش دهند. چنین رویه‌ای می‌تواند به انتقال ناعادلانه ریسک به کاربران منجر شود. اتحادیه اروپا برای مقابله با این وضعیت، مقررات الزام‌آور متعددی را تصویب کرده است. قانون حمایت از داده‌ها (GDPR 2016/679) برای نخستین بار حق دسترسی کاربران به منطق تصمیم‌گیری خودکار را پیش‌بینی کرد. قانون جدید هوش مصنوعی (AI Act 2024) نیز برای سامانه‌های پر ریسک، تعهدات سخت‌گیرانه شفافیت و نظارت انسانی وضع کرده است. این قوانین به دنبال ایجاد تعادل میان حمایت از نوآوری و تضمین حقوق بنیادین کاربران هستند.

با وجود تصویب مقررات ملی و منطقه‌ای، همچنان پرسش‌های اساسی درباره نحوه انتساب مسئولیت به تصمیمات هوش مصنوعی بی‌پاسخ مانده است. پیچیدگی فنی الگوریتم‌ها باعث شده است که رابطه میان فعل زیان‌بار و عامل انسانی تضعیف شود. در بسیاری موارد نمی‌توان تعیین کرد که خطا ناشی از

طراحی، داده، یا به‌روزرسانی خودکار سامانه بوده است. این وضعیت موجب چالش‌های جدی در دستیابی به عدالت جبرانی می‌شود. دادگاه‌ها نیز در نبود چارچوب روشن، با ابهام در تشخیص تقصیر یا قصور مواجه‌اند. چنین ابهامی امنیت حقوقی جامعه را نیز تحت تأثیر قرار می‌دهد. از این‌رو نیاز به تدوین قواعد شفاف و چندلایه بیش از پیش احساس می‌شود. این قواعد باید تمام بازیگران از جمله توسعه‌دهندگان، بهره‌برداران و کاربران را پوشش دهد. بر همین اساس، پرسش اصلی پژوهش‌ناظر بر تعیین مسئولیت حقوقی تصمیمات الگوریتمی و جایگاه کاربران در این فرآیند است.

پژوهش علی‌پناهی و همکاران (۱۴۰۳: ۸۷) با بررسی مقررات مسئولیت مدنی در اتحادیه اروپا نشان می‌دهد که تعیین چارچوب‌های روشن مسئولیت می‌تواند تضمین‌کننده جبران خسارت قربانیان در حوزه هوش مصنوعی باشد. آنان تأکید می‌کنند که در موارد ورود زیان، اعمال مسئولیت محض بر توسعه‌دهندگان و اپراتورها مانع از فرار مسئولیت خواهد شد. این دیدگاه بیانگر آن است که اتکای صرف به نظریه تقصیر برای تحلیل تصمیمات الگوریتمی کافی نیست. در همین راستا، ذاکری‌نیا (۱۴۰۲: ۷۴) بر ناکارآمدی معیارهای سنتی نظیر تقصیر و رابطه سببیت تأکید کرده و ضرورت بازاندیشی بنیادی در این مفاهیم را مطرح می‌کند. وی پیشنهاد می‌دهد که ترکیب مسئولیت محض و سنتی می‌تواند پاسخی واقع‌بینانه به پیچیدگی الگوریتم‌ها باشد. تخشید (۱۴۰۰: ۶۱) نیز با بررسی مسأله شخصیت حقوقی سامانه‌های هوشمند، پرسش امکان شناسایی این سامانه‌ها به عنوان فاعل مستقل را مطرح می‌کند. اما ولی‌پور و اسماعیلی (۱۴۰۰: ۹۴) بر فقدان شخصیت حقوقی هوش مصنوعی تأکید کرده و معتقدند مسئولیت باید بر اشخاص انسانی متمرکز بماند. از سوی دیگر، حکمت‌نیا و همکاران (۱۳۹۸: ۱۳۳) نشان می‌دهند که خودمختاری ربات‌های هوشمند، مسئولیت صرف کاربران را ناکافی می‌سازد. این یافته‌ها ضرورت طراحی چارچوب‌های جایگزین برای تضمین حمایت مؤثر از زیان‌دیده را تقویت می‌کند. گارون (۲۰۲۴: ۲۴) نیز با بررسی خلأهای موجود در انتساب مستقیم خروجی الگوریتم به فاعل انسانی، امکان اعطای شخصیت حقوقی محدود به سامانه‌های هوشمند را مطرح می‌سازد. وی در عین حال بر لزوم بازنگری در معیار قابلیت پیش‌بینی تأکید می‌کند. در نهایت، مقاله او مدل ترکیبی مسئولیت محض و تقصیر را کارآمدترین الگو می‌داند. همچنین،

جادج و همکاران (۲۰۲۴) با تأکید بر ناآگاهی کاربران از پیامدهای حقوقی تصمیمات خودکار، مسئولیت را امری زنجیره‌ای میان طراحان، توسعه‌دهندگان و بهره‌برداران می‌دانند. آنان شفافیت و توضیح‌پذیری را پیش شرط تحقق عدالت حقوقی معرفی می‌کنند. بررسی آنان از تجربه اتحادیه اروپا نیز ضرورت تدوین مقررات الزام‌آور ملی و فراملی را برجسته می‌سازد.

۱- چالش‌های بنیادین مسئولیت مدنی در مواجهه با تصمیمات الگوریتمی

ظهور سامانه‌های هوش مصنوعی با ساختارهای خودآموز و غیرخطی، بنیان‌های کلاسیک مسئولیت مدنی را با چالشی جدی مواجه کرده است. معیارهای کلاسیک تقصیر، قابلیت پیش‌بینی و رابطه سببیت، هنگامی که تصمیم به‌جای اراده انسانی از سازوکار آماری - الگوریتمی ناشی می‌شود، کارایی روش‌شناختی خود را از دست می‌دهند. از این‌رو، تحلیل مبانی نظری مسئولیت نیازمند بازخوانی غایت‌شناسانه وظایف مراقبتی در طراحی، آموزش و بهره‌برداری از الگوریتم‌هاست. پژوهش‌های پژوهشی داخلی بر ضرورت تفکیک میان سامانه‌های صرفاً خودکار و سامانه‌های یادگیرنده و تأثیر این تفکیک بر قلمرو انتساب مسئولیت تأکید نهاده‌اند (اکبری و همکاران، ۱۴۰۲: ۴۵).

بدین‌منظور، باید «استاندارد توجه متعارف» از قالب رفتار انسانی به «استاندارد دقت الگوریتمی» منتقل شود تا قابلیت اعمال در بستر تصمیم‌گیری‌های داده‌محور فراهم آید. همچنین، اصل حمایت از بزه‌دیده اقتضا دارد که نظام دادرسی در مواجهه با عدم توضیح‌پذیری خروجی‌ها از قواعد اثباتی انعطاف‌پذیر بهره گیرد. در این چارچوب، فرض‌های قابل نقض دربارهٔ انحراف الگوریتم از پارامترهای اعلامی می‌تواند نقش جایگزین در بار اثبات ایفا کند. پژوهش‌های داخلی اخیر با تبیین مرز میان ابزار و فاعل در سامانه‌های یادگیرنده، پیامدهای آن را بر بازتعریف قلمرو «فاعل حقوقی» و «ابزار خطرناک» نشان داده‌اند (صابری و همکاران، ۱۴۰۳: ۸۲). نتیجه آنکه، مبانی مسئولیت مدنی در پرتو هوش مصنوعی مستلزم انتقال از الگوی «فاعل - فعل» به الگوی «نظام - مبتنی بر خطر» است. چنین انتقالی تنها با پالایش مفهومی عناصر مسئولیت و بازمهندسی قواعد اثباتی میسر خواهد شد.

در سطح تطبیقی، پژوهش‌های بین‌المللی نشان می‌دهد که رکن «توضیح‌پذیری» به‌عنوان شرط مشروعیت تصمیمات الگوریتمی،

نقشی بنیادین در بازتعریف عناصر مسئولیت یافته است. در فقدان شفافیت قابل سنجش، هم معیار پیش‌بینی‌پذیری مخدوش می‌شود و هم امکان دفاع مؤثر برای اصحاب دعوا محدود می‌گردد. از این رهگذر، نظریه‌های نوین، تکلیف «افشای فنی معنادار» و «ممیزی الگوریتمی» را در شمار تعهدات پیش‌قراردادی و پس‌قراردادی قرار می‌دهند. پژوهش‌های تخصصی در حوزه حقوق فناوری با تأکید بر خطر تبعیض و خطای سیستمی، تنظیم‌گری مبتنی بر خطر و حق اعتراض مؤثر به تصمیمات خودکار را لازمه رعایت اصول دادرسی منصفانه دانسته‌اند (Veale & Zuiderveen Borgesius, 2021: 410). چنین تلقی‌ای بر آن است که توزیع بار خطر نباید صرفاً تابع مالکیت نرم‌افزار باشد، بلکه باید به میزان کنترل مؤثر بر چرخه عمر الگوریتم گره بخورد. در نتیجه، «زنجیره مسئولیت» از مرحله طراحی تا استقرار و پایش مستمر تعریف می‌شود و هر حلقه تعهدات متمایز می‌یابد. این رویکرد، توجیه‌های جاری انتقال از مسئولیت فردی به مسئولیت‌های نهادی و شبکه‌ای را فراهم می‌سازد. افزون بر آن، توصیه‌های سیاستی نهادهای اروپایی بر لزوم استانداردگذاری برای ثبت ردپای تصمیم و ایجاد واجدان صلاحیت برای حسابرسی مستقل تأکید دارند (Mökander et al, 2022: 252-253). برآیند این پژوهش‌های، پیوند وثیق میان مشروعیت حقوقی و قابلیت حسابرسی فنی را صورت‌بندی می‌کند. بدین‌سان، مبانی نظری مسئولیت ناگزیر به ادغام معیارهای حقوقی و سنجش‌های فنی است.

مسأله «سببیت» در بستر تصمیمات الگوریتمی، به‌واسطه هم‌پوشانی داده‌ها، مدل‌ها و مداخله‌های مکرر به‌روزرسانی، از قالب خطی و تک‌عاملی خارج شده و مستلزم الگوهای توصیفی پیچیده‌تر است. در فقدان دسترسی به مسیرهای درونی مدل، تحلیل نسبت علیت میان داده ورودی و زیان خروجی به دشواری قابل احراز است. بر این مبنا، بازآرایی قواعد انتساب از طریق پذیرش «سببیت نهادی» و «سببیت جمعی زنجیره‌ای» قابل دفاع می‌نماید. پژوهش‌های داخلی با تأکید بر لزوم انتقال تمرکز از «کنش فردی» به «معماری حاکمیت داده و مدل»، راهبردهای اثباتی بدیل از قبیل امارات ناشی از عدم انطباق با استانداردهای اعلامی را پیشنهاد کرده است (علی‌نقیان و همکاران، ۱۴۰۲: ۶۳). در این الگو، تخطی از تعهدات شفافیت، ممیزی و ثبت سوابق فنی، به‌منزله قرینه‌ای بر دخالت سببیت حقوقی تلقی می‌شود. هم‌زمان، نظریه «قصور در نظارت

در ساحت حقوق داخلی، «تنظیم‌گری مبتنی بر ریسک» اقتضا می‌کند که درجات تعهد و مسئولیت به بر اساس نوع کاربرد سامانه، حساسیت حوزه و شدت پیامدهای بالقوه، طبقه‌بندی شود. پیامد هنجاری این طبقه‌بندی آن است که در حوزه‌های پر مخاطره، بار مراقبت و شفافیت تشدید و در حوزه‌های کم‌خطر، انعطاف مقرراتی حفظ گردد. این تمایزگذاری باید با اصول برابری، منع تبعیض و تناسب همخوان باشد تا از تحمیل هزینه‌های ناموجه جلوگیری شود. پژوهش‌های داخلی جدید، ضمن تأکید بر لزوم ارزیابی اثرات الگوریتمی پیشینی، «حق آگاهی از منطق تصمیم» و «حق اعتراض مؤثر» را جزء لاینفک حمایت‌های مصرف‌کننده و حریم خصوصی دانسته‌اند (حاجی‌اسماعیلی، ۱۴۰۳: ۱۰۲).

در پرتو این حقوق، عدم افشای مقتضی یا امتناع از ارائه تبیین قابل سنجش، می‌تواند به‌عنوان نقض تعهدات قانونی و موجد مسئولیت تلقی گردد. افزون بر این، پذیرش الگوی «مسئولیت لایه‌ای» امکان توزیع ریسک میان تولیدکننده، ارائه‌دهنده خدمت و بهره‌بردار را به‌نحوی عادلانه‌تر فراهم می‌کند. این الگو، همسو با اصل پیشگیری، انگیزه‌های مراقبتی را در تمام طول زنجیره تقویت می‌نماید. در همین راستا، حمایت‌های بیمه‌ای و سازوکارهای جبرانی مکمل نیز باید به‌طور ساختاری در طراحی حکمرانی الگوریتمی ادغام شوند. نتایج هنجاری چنین ادغامی، افزایش کارایی جبرانی و ارتقاء اعتماد مشروع به تصمیمات فناورانه است (بحرکاظمی و محمودی، ۱۴۰۳: ۵۹). ماهیت فرامرزی سامانه‌های هوش مصنوعی چالش‌های پیچیده‌ای در حوزه تعارض قوانین و صلاحیت قضایی ایجاد کرده است. پژوهش‌های زیان‌بار ناشی از تصمیمات الگوریتمی ممکن است در قلمروهای گوناگون جغرافیایی واقع شود، درحالی‌که سازوکارهای سنتی تعیین صلاحیت مبتنی بر مکان وقوع فعل زیان‌بار، پاسخگوی این وضعیت نیست. از این‌رو، بازاندیشی در معیارهای صلاحیت، از جمله الحاق مکان بروز اثر و محل استقرار زیرساخت داده به قلمرو صلاحیت، ضرورت دارد. پژوهش‌های داخلی با تأکید بر «اثرگذاری برون‌مرزی تصمیمات خودکار» لزوم بازنگری در قواعد آیین دادرسی مدنی را یادآور شده‌اند. این تحول نه‌تنها جنبه‌ی رویه‌ای دارد بلکه به‌طور مستقیم بر حقوق مادی اصحاب دعوا نیز اثر می‌گذارد. در غیاب هماهنگی مقرراتی، خطر «فرار تنظیمی» و جابه‌جایی شرکت‌ها به حوزه‌های قضایی با

الگوریتمی» به‌عنوان صورت‌بندی نوین تقصیر، قلمرو مراقبت حرفه‌ای را از مصرف به تولید و نگاه‌داری مدل توسعه می‌دهد. چنین صورت‌بندی‌ای می‌تواند مبانی مسئولیت قهری را در فقدان دسترسی به منطق درونی تصمیم تقویت کند. علاوه بر این، پژوهش‌های داخلی اخیر بر ضرورت طراحی نظام‌های بیمه‌ای مکمل برای پوشش خطرهای غیرقابل انتساب فردی تأکید نهاده‌اند (سیفی و رزمخواه، ۱۴۰۱: ۸۸). بدین‌سان، سببیت از پدیده‌ای تک‌عاملی به سازوکار نهادی چندسطحی تبدیل می‌شود. این تحول، کارایی جبرانی و پیشگیری عمومی را در حوزه مسئولیت مدنی ارتقاء می‌بخشد.

۲- تحول نظام‌های حقوقی در مواجهه با ریسک‌های هوش مصنوعی

در سوبه تطبیقی، مبنای مشروعیت انتساب مسئولیت به سامانه‌های پر ریسک بر محور «حکمرانی ریسک» و «پاسخگویی الگوریتمی» بازتعریف می‌شود. در این برداشت، معیارهای سنتی مبتنی بر اراده کنشگر انسانی جای خود را به معیارهای ساختاری مبتنی بر کنترل مؤثر، قابلیت ممیزی و ردیابی می‌دهند. پیامد هنجاری آن است که عدم تأمین زیرساخت‌های پاسخگویی، خود صورت تقصیر ساختاری تلقی می‌شود. پژوهش‌های تخصصی معاصر درباره نظم حقوقی اروپا نشان داده است که اصول اخلاقی ناظر بر شفافیت و مسئولیت‌پذیری، از سطح توصیه به سطح قواعد الزام‌آور ارتقاء یافته و به زبان حقوقی قابل اجرا ترجمه شده‌اند (Burri, 2023: 74). در این چارچوب، تعهدات مثبت برای مستندسازی چرخه حیات مدل، آزمون سوگیری، و تضمین حق بازبینی انسانی از ارکان مسئولیت قلمداد می‌گردند. همچنین، پذیرش «استاندارد الگوریتم معقول» به‌مثابه‌ی بدل «انسان متعارف» در برخی پیشنهادها، صورت‌بندی جدیدی از معیار مراقبت را پدید آورده است. چنین استاندارد، سنجش رفتار الگوریتم را با معیارهای فنی روز و عرف حرفه‌ای همساز می‌کند. افزون بر آن، پژوهش‌های تحلیلی حقوق و فناوری بر نقش تحلیل‌پذیری فنی در تضمین حق دادرسی منصفانه و برابری سلاح‌ها تأکید دارد (Ashley, 2019: 128). نتیجه، پیوند مفهومی میان کارآمدی جبرانی و قابلیت فنی تبیین است. بدین‌سان، مبانی نظری مسئولیت، صورت‌بندی چرخه‌ای میان حق و فناوری را به نمایش می‌گذارد.

قواعد سهل‌گیرانه افزایش می‌یابد. بدین ترتیب، ضرورت تنظیم فراملی قواعد مربوط به حکمرانی داده و الگوریتم آشکار می‌شود. افزون بر آن، پیشنهاد شده است که نظام‌های «شناخت متقابل» در صدور گواهی‌نامه‌های اعتماد الگوریتمی به‌عنوان ابزار هماهنگ‌سازی در سطح منطقه‌ای به کار گرفته شوند (Shneiderman, 2022: 93). نتیجه این مباحث، تقویت کارایی قضایی و پیشگیری از تضییع حقوق زیان‌دیده در سطح فراملی است.

۳- پیامدهای حقوقی نوظهور هوش مصنوعی

یکی از مسائل بنیادین، پرسش از امکان اعطای شخصیت حقوقی مستقل به سامانه‌های هوش مصنوعی است. چنین بحثی، در پیوند مستقیم با قابلیت انتساب حقوق و تکالیف به این سامانه‌ها قرار دارد. پذیرش شخصیت حقوقی برای هوش مصنوعی مستلزم بازتعریف مفهوم «فاعل حقوقی» و تفکیک آن از ابزار صرف است. پژوهش‌های داخلی این رویکرد را عمدتاً با تردید تلقی کرده و بر ضرورت حفظ چارچوب‌های سنتی عاملیت حقوقی تأکید نموده است (موسوی، ۱۴۰۲: ۱۲-۱۴). استدلال اصلی این است که شخصیت حقوقی بدون توانایی ادراک تکلیف و اراده آگاهانه فاقد مبنای هنجاری است. در مقابل، برخی نظریات بر امکان پذیرش شخصیت حقوقی محدود و ابزاری برای سامانه‌های خاص استدلال کرده‌اند. این دیدگاه بر معیار کارکردی ایفای نقش در تعاملات اقتصادی و اجتماعی تکیه دارد. بااین‌حال، خطر تضعیف اصل پاسخگویی و انتقال ناموجه ریسک همچنان باقی است. مطالعات تطبیقی نیز نشان می‌دهد که پذیرش شخصیت حقوقی برای هوش مصنوعی باید با محدودیت‌های مضیق و نظارت قضایی سخت‌گیرانه همراه باشد (Solum, 2020: 158). در نهایت، به نظر می‌رسد نظام حقوقی فعلاً باید به سمت توسعه الگوهای مسئولیت مدنی و کیفری بر محور کنترل‌کنندگان و بهره‌برداران حرکت کند.

از منظر حقوق اقتصادی، تأثیر تصمیمات الگوریتمی بر روابط قراردادی و الزامات ناشی از اصل حسن‌نیت نیز قابل تأمل است. هوش مصنوعی می‌تواند به‌عنوان عامل اجرایی قرارداد عمل کرده و تصمیماتی اتخاذ کند که پیامدهای مستقیم بر تعادل قراردادی دارند. این وضعیت، ضرورت بازتعریف مفاهیم سنتی مانند «رضایت»، «اراده» و «عیب رضا» را آشکار می‌سازد. پژوهش‌های داخلی تأکید دارند که در قراردادهایی که اجرای آن‌ها بر عهده سامانه‌های هوشمند قرار می‌گیرد، باید سطح

خاصی از شفافیت و قابلیت پیش‌بینی برای طرفین تضمین شود. در غیر این صورت، اصل انصاف قراردادی و حمایت از طرف ضعیف نقض خواهد شد. افزون بر این، تعهدات اطلاع‌رسانی و افشای ویژگی‌های سامانه بخشی از الزامات پیش‌قراردادی تلقی می‌گردد. این تعهدات با اصول حاکم بر حمایت از مصرف‌کننده پیوند می‌خورند و تخلف از آن‌ها می‌تواند موجب بطلان یا قابلیت فسخ قرارداد شود. در سطح بین‌المللی نیز مطالعات نشان داده‌اند که تصمیمات الگوریتمی می‌توانند تعادل قراردادی را مختل کرده و زمینه‌ساز مسئولیت ناشی از نقض حسن‌نیت شوند (Edwards & Veale, 2017: 42).

بنابراین، نظام حقوقی باید قواعدی وضع کند که همسو با اصول انصاف و کارایی اقتصادی، از بروز بی‌عدالتی جلوگیری نماید. نهایتاً، پیوند میان هوش مصنوعی و حقوق کیفری پرسش‌های پیچیده‌ای را درباره حدود مسئولیت و قابلیت انتساب جرم ایجاد می‌کند. چالش اصلی این است که اگر الگوریتم منجر به ارتکاب رفتار مجرمانه شود، تعیین فاعل جرم چگونه امکان‌پذیر خواهد بود. اصول سنتی مسئولیت کیفری بر آگاهی و قصد فاعل انسانی استوارند، درحالی‌که الگوریتم فاقد اراده و شعور مستقل است. پژوهش‌های داخلی بر این نکته تأکید کرده است که در چنین مواردی باید مسئولیت متوجه اشخاصی باشد که بر طراحی، آموزش و بهره‌برداری از الگوریتم نظارت داشته‌اند (آقا امینی فشمی و همکاران، ۱۴۰۲: ۶۹). این دیدگاه مبتنی بر اصول سیاست جنایی پیشگیرانه است و هدف آن جلوگیری از گریز اشخاص حقیقی یا حقوقی از پاسخگویی کیفری است. در حوزه تطبیقی، برخی اندیشمندان استدلال کرده‌اند که می‌توان الگوی «قصور ساختاری» را به‌عنوان مبنای انتساب مسئولیت کیفری در نظر گرفت، بدین معنا که عدم اتخاذ تدابیر کافی در طراحی و نظارت می‌تواند معادل تقصیر جزایی تلقی گردد (Calo, 2015: 39-41). چنین تحولی، هم‌سو با اصل پیشگیری عمومی، به ارتقاء امنیت اجتماعی در مواجهه با فناوری‌های نوین کمک خواهد کرد.

۴- مسئولیت مدنی ناشی از تصمیم‌گیری‌های الگوریتمی

مسئولیت مدنی در نظام حقوقی سنتی بر پایه‌ی عناصر تقصیر، رابطه سببیت و وقوع ضرر استوار بوده است و این عناصر نقش بنیادین در تضمین جبران خسارت ایفا کرده‌اند. در حوزه تصمیم‌گیری‌های الگوریتمی، به‌واسطه ماهیت غیرشفاف و پیچیده سامانه‌های هوش مصنوعی، تطبیق این عناصر با واقعیت

بی‌توجهی به رعایت استانداردهای طراحی یا آموزش الگوریتم می‌تواند به‌عنوان قرینه‌ای بر تحقق سببیت تلقی شود (رجبی، ۱۳۹۸: ۴۶۲-۴۶۵). بنابراین، به‌جای اتکا به الگوهای صرفاً جمعی، مناسب‌تر آن است که نظام حقوقی با بهره‌گیری از معیارهای متعارف و کارکردی سببیت، ضمن توجه به اقتضائات فناوری‌های نوین، همچنان بر ضرورت احراز رابطه حقیقی میان رفتار و ضرر تأکید کند تا هم حمایت از زیان‌دیده حفظ شود و هم از ایجاد مسئولیت غیرمنصفانه جلوگیری گردد.

یکی دیگر از مباحث اساسی در مسئولیت مدنی ناشی از تصمیم‌گیری‌های الگوریتمی، مسأله پیش‌بینی‌پذیری است که به‌عنوان شاخصی هنجاری در احراز تعهدات مراقبتی شناخته می‌شود. در حقوق سنتی، معیار پیش‌بینی بر پایه رفتار انسان متعارف بنا نهاده شده و از کنشگر انتظار می‌رود پیامدهای معقول فعل خود را لحاظ کند. در مقابل، الگوریتم‌های یادگیرنده قادر به تولید خروجی‌هایی هستند که حتی طراحان آن‌ها نیز پیش‌بینی دقیق آن‌ها را ندارند. این امر منجر به چالش اساسی در تعیین حد پیش‌بینی‌پذیری می‌شود و می‌تواند زیان‌دیدگان را در وضعیت آسیب‌پذیری مضاعف قرار دهد. از این‌رو، پیشنهاد شده است که به‌جای معیار مطلق پیش‌بینی، معیار «پیش‌بینی معقول» مبتنی بر رعایت استانداردهای فنی پذیرفته شود. این معیار متکی بر این پرسش است که آیا طراح یا بهره‌بردار اقدامات متعارف برای پیشگیری از زیان را انجام داده است یا خیر. در این صورت، مسئولیت مدنی بر مبنای عدم رعایت استانداردهای حرفه‌ای و نه پیش‌بینی مطلق استوار می‌گردد. چنین صورتی از مسئولیت می‌تواند ضمن حمایت از زیان‌دیده، مانع از توقف نوآوری نیز گردد. در نهایت، پذیرش این معیار سازوکاری برای توازن میان پیشرفت فناوریانه و تضمین حقوق زیان‌دیده به وجود خواهد آورد (Chesterman, 2020: 824-826). بنابراین، عنصر پیش‌بینی‌پذیری در حقوق مدنی باید همگام با تحولات فناوریانه بازتعریف شود.

یکی از جلوه‌های نوین در مسئولیت مدنی ناشی از تصمیم‌گیری‌های الگوریتمی، مسأله تعدد بازیگران و لزوم شناسایی مسئولیت تضامنی است. زنجیره تولید و بهره‌برداری از هوش مصنوعی شامل طراحان نرم‌افزار، ارائه‌دهندگان خدمات پردازش داده و کاربران نهایی است که هر یک در وقوع زیان نقش مستقیم یا غیرمستقیم ایفا می‌نمایند. در چنین وضعیتی، تقسیم مسئولیت بر مبنای معیار سنتی فردی نمی‌تواند پاسخگوی

عینی با دشواری جدی روبه‌رو شده است. عنصر تقصیر در معنای کلاسیک خود متکی بر رفتار غیرمتعارف شخص انسانی است، درحالی‌که الگوریتم فاقد اراده انسانی بوده و از قواعد آماری تبعیت می‌کند. همین امر موجب خلأ در امکان انتساب مستقیم تقصیر به سامانه می‌شود و نیاز به بازخوانی مفهومی این عنصر را ایجاب می‌نماید. در این راستا، برخی از نظریه‌ها بر پذیرش «قصور در نظارت الگوریتمی» به‌عنوان صورت نوین تقصیر تأکید دارند که در آن، عدم کنترل کافی بر فرایندهای فنی مصداق بی‌مبالاتی تلقی می‌گردد. چنین بازتعریفی به حفظ کارایی نهاد مسئولیت مدنی در برابر فناوری‌های نوین یاری می‌رساند. افزون بر این، این صورت‌بندی به زیان‌دیده امکان می‌دهد بار اثباتی خود را از رفتار فنی به تعهدات نظارتی منتقل نمایند. از منظر عدالت ترمیمی، این تحول ضمن تقویت حق دسترسی به جبران، توازن میان توسعه فناوری و حمایت از افراد را برقرار می‌کند (کریمی و علی‌آبادی، ۱۴۰۳: ۱۳۱-۱۲۹). در نتیجه، بازتعریف عنصر تقصیر در پرتو تصمیم‌گیری‌های الگوریتمی ضرورتی گریزناپذیر برای حفظ مشروعیت نهاد مسئولیت مدنی است.

عنصر رابطه سببیت در مسئولیت مدنی سنتی به‌عنوان پیوندی میان فعل زیان‌بار و ضرر شناخته می‌شود و بار اصلی اثبات آن بر دوش مدعی قرار دارد. در بستر تصمیم‌گیری‌های الگوریتمی، به دلیل پیچیدگی فرایندهای یادگیری ماشین و تغییرپذیری مداوم مدل‌ها، احراز این رابطه با دشواری جدی مواجه است و مسیر تحقق ضرر به‌صورت خطی و قابل بازسازی نیست. این وضعیت امکان اقامه دعوی مؤثر را برای زیان‌دیده محدود کرده و کارکرد جبرانی نهاد مسئولیت را تضعیف می‌کند. از این‌رو، برخی حقوقدانان نظریه «سببیت نهادی» یا «سببیت جمعی» را پیشنهاد کرده‌اند تا مسئولیت میان مجموعه‌ای از عوامل زنجیره تولید و بهره‌برداری توزیع گردد. هرچند این رویکرد می‌تواند بار اثبات را از دوش زیان‌دیده بردارد، اما این نظریه با ضعف بنیادی مواجه است؛ زیرا ممکن است یکی از عوامل، نقشی واقعی در ورود خسارت نداشته باشد اما صرفاً به‌دلیل انتساب جمعی مسئول شناخته شود که این امر هم خلاف اصل شخصی‌بودن مسئولیت و هم موجب تحمیل بار ناموجه بر عامل بی‌تقصیر است. دکتر یزدانیان نیز با تأکید بر ضرورت احراز رابطه واقعی میان فعل و ضرر، هشدار می‌دهد که تضعیف معیار سنتی سببیت می‌تواند به شکل‌گیری مسئولیت‌های غیرعادلانه و گسترش بی‌ضابطه حوزه مسئولیت بینجامد (یزدانیان، ۱۴۰۳: ۴۶-۴۰). در کنار این تحلیل،

بدین‌سان، تحقق عدالت ترمیمی تنها از رهگذر همگرایی این دو سازوکار امکان‌پذیر خواهد بود.

۵- جایگاه حقوقی کاربران در فرآیند تصمیم‌گیری هوش

مصنوعی

کاربر به‌عنوان مخاطب نهایی تصمیمات الگوریتمی، دارای جایگاهی دوگانه در نظام حقوقی است: از یک‌سو بهره‌بردار خدمات و از سوی دیگر، شخصی که می‌تواند در معرض زیان ناشی از تصمیم خودکار قرار گیرد. این دوگانگی سبب شده است که شناسایی دقیق حقوق و تکالیف کاربران در فرآیند تصمیم‌گیری هوش مصنوعی اهمیت بنیادین بیابد. در حقوق سنتی، کاربر عمدتاً در مقام مصرف‌کننده تلقی شده و حقوق وی در چارچوب قواعد حمایت از مصرف‌کننده تعریف می‌شده است. اما در بستر تصمیمات الگوریتمی، حقوق کاربر تنها به بعد قراردادی محدود نمی‌شود بلکه ابعاد اساسی چون حریم خصوصی، حق بر آگاهی و حق اعتراض را نیز شامل می‌گردد. از این منظر، کاربر نه صرفاً طرف قرارداد بلکه ذی‌نفعی در روند تولید و اجرای تصمیم شناخته می‌شود. همین امر لزوم ارتقاء جایگاه وی به سطحی فراتر از مصرف‌کننده صرف را توجیه می‌کند. پژوهش‌های داخلی در این خصوص تصریح کرده‌اند که «حق دسترسی به منطق تصمیم» و «حق بازبینی انسانی» باید به‌عنوان حقوق بنیادین کاربر در برابر سامانه‌های هوشمند شناسایی شود (رضوی، ۱۴۰۱: ۸۴). این تحول از حیث تقنینی مستلزم وضع قواعد الزام‌آور و نه صرفاً توصیه‌های اخلاقی است. در نتیجه، جایگاه حقوقی کاربر در این حوزه باید به‌مثابه بخشی از حقوق بنیادین شهروندی در فضای دیجیتال تلقی گردد.

از حیث تعهدات قراردادی، کاربران در بسیاری موارد بدون آنکه به‌طور واقعی شرایط استفاده از خدمات الگوریتمی را درک کنند، با آن موافقت می‌نمایند. این امر ناشی از پیچیدگی اسناد حقوقی، زبان فنی قراردادها و سرعت دسترسی به خدمات است که منجر به تضعیف اصل رضایت آگاهانه می‌شود. در این زمینه، ضرورت دارد که قواعد سنتی مربوط به «عیب رضا» بازخوانی و متناسب با اقتضائات فناوری‌های نوین بازتعریف شود. در غیر این صورت، کاربر در موقعیتی نابرابر نسبت به ارائه‌دهنده خدمات قرار گرفته و اصل توازن قراردادی نقض خواهد شد. علاوه بر این، تعهدات اطلاع‌رسانی و شفافیت باید به‌عنوان بخشی از تکالیف پیش‌قراردادی بر عهده ارائه‌دهنده خدمات نهاده شود. در این معنا، رضایت آگاهانه کاربر به‌عنوان شرط صحت قراردادهای

واقعیت‌های فنی باشد. نظریه مسئولیت تضامنی بر آن است که زیان‌دیده بتواند برای جبران خسارت به هر یک از بازیگران مراجعه نماید بدون آنکه ناگزیر به تعیین دقیق سهم تقصیر هر فرد باشد. این رویکرد ضمن تسهیل دسترسی بزه‌دیده به جبران، انگیزه‌های مراقبتی را در تمامی حلقه‌های زنجیره افزایش می‌دهد. در این چارچوب، امکان رجوع متقابل میان مسئولان برای باز توزیع بار زیان نیز پیش‌بینی می‌شود تا عدالت نسبی در روابط داخلی آنان برقرار گردد. به همین دلیل، برخی پژوهش‌های حقوقی پیشنهاد کرده‌اند که الگوی مسئولیت تضامنی باید به‌صورت قاعده‌ای عام در حوزه الگوریتم‌های پر ریسک پذیرفته شود (Tai, 2022: 121). چنین پذیرشی می‌تواند مانع از فروپاشی حمایت‌های جبرانی در نظام مدنی شود. علاوه بر آن، با تقویت ضمانت اجرای تضامنی، اصول پیشگیری عمومی نیز به‌صورت مؤثرتر تحقق می‌یابند. بدین ترتیب، مسئولیت تضامنی در این حوزه ابزار کارآمدی برای تحقق هم‌زمان عدالت ترمیمی و پیشگیری از وقوع زیان خواهد بود.

بعد دیگر مسئولیت مدنی در تصمیم‌گیری‌های الگوریتمی، ضرورت طراحی سازوکارهای بیمه‌ای و جبرانی مکمل است که در کنار قواعد سنتی به تضمین کارایی جبرانی کمک نمایند. در غیاب چنین ابزارهایی، بار مالی سنگین زیان‌ها ممکن است موجب ناکارآمدی کلی نهاد مسئولیت گردد و حقوق زیان‌دیده بالاجبران باقی بماند. بیمه اجباری در حوزه‌های پرخطر فناورانه می‌تواند نقش ابزاری اساسی در توزیع ریسک اجتماعی ایفا کند. بدین معنا که خسارت‌های ناشی از خطا یا انحراف الگوریتم میان جامعه و بهره‌برداران توزیع می‌گردد و از تمرکز زیان بر یک شخص جلوگیری می‌شود. علاوه بر آن، نهادهای بیمه‌ای می‌توانند با وضع شروط قراردادی، بهره‌برداران را به رعایت استانداردهای فنی و نظارتی ملزم سازند. این امر موجب ارتقاء سطح مراقبت و پیشگیری از بروز خطاهای الگوریتمی خواهد شد. در پژوهش‌های حقوق داخلی نیز تأکید شده است که طراحی صندوق‌های جبرانی مکمل می‌تواند به‌عنوان ضمانت اجرای ثانویه برای جبران خسارت‌های غیرقابل‌پیش‌بینی عمل کند (مهدوی، ۱۳۹۸: ۵۳). چنین راهکاری هم‌سو با اصول عدالت توزیعی و کارایی اقتصادی تلقی می‌شود. در نهایت، ترکیب قواعد مسئولیت مدنی و نهادهای بیمه‌ای می‌تواند بنیانی پایدار برای حمایت از زیان‌دیده در حوزه هوش مصنوعی فراهم سازد.

این حقوق نیازمند شفافیت عملیاتی و ثبت ردپای تصمیمات الگوریتمی است. در نتیجه، جایگاه کاربر از یک مصرف‌کننده منفعل به یک صاحب حق فعال تبدیل می‌شود. این تحول موجب استقرار توازن جدیدی میان نوآوری فناوریانه و صیانت از حریم خصوصی می‌گردد. در نهایت، تقویت این نظام‌ها شرط ضروری مشروعیت تصمیم‌گیری الگوریتمی در دولت‌های معاصر است.

در حوزه نظارت بر تصمیمات خودکار، حق اعتراض و حق بازبینی انسانی به‌عنوان دو رکن بنیادین تضمین‌کننده دادرسی منصفانه در فضای دیجیتال شناخته می‌شوند. این حقوق اجازه می‌دهند که کاربران در برابر تصمیماتی که بر زندگی، اشتغال، سلامت یا وضعیت حقوقی آنان اثرگذار است، از سازوکارهای کنترلی مؤثر برخوردار باشند. نظام‌های حقوقی پیشرفته تصریح می‌کنند که هیچ تصمیمی نباید صرفاً مبتنی بر پردازش خودکار و بدون نظارت انسانی اتخاذ شود. از این منظر، الزام به توضیح‌پذیری و ارائه دلایل تصمیم برای کاربر، بخشی از حقوق دفاعی او تلقی می‌شود. در این میان، ایجاد سازوکارهای شکایت و امکان مراجعه به مراجع مستقل ضروری است تا کاربران بتوانند از تضییع حقوق خود جلوگیری کنند. بسیاری از کشورها، از جمله اعضای اتحادیه اروپا، مدل‌های چندلایه نظارتی را برای کنترل سامانه‌های پرخطر توسعه داده‌اند. ایران نیز نیازمند توسعه ساختارهای حقوقی مشخص برای بازبینی تصمیمات خودکار در بخش‌های عمومی و خصوصی است. برخی پژوهشگران معتقدند که غیاب این سازوکارها موجب شکل‌گیری نابرابری حقوقی و آسیب‌پذیری کاربران می‌شود (صابری، ۱۴۰۳: ۱۰۷).

افزون بر این، تضمین امنیت داده و پیشگیری از نفوذ یا دست‌کاری الگوریتم‌ها از الزامات بنیادین حکمرانی دیجیتال است. این الزام ایجاب می‌کند که لایه‌های امنیتی مانند رمزنگاری، کنترل دسترسی و ثبت سوابق فنی در تمامی مراحل چرخه تصمیم‌گیری فعال باشند. از سوی دیگر، حقوق صاحبان داده نباید در هیچ شرایطی نقض شود و نظام تنظیم‌گری باید مانع انتقال ناعادلانه ریسک به کاربران گردد. این نظام همچنین باید امکان رفع خطا، توقف تصمیم‌گیری خودکار و دخالت انسان را فراهم کند. اجرای این اصول موجب ارتقاء اعتماد مشروع شهروندان به استفاده از فناوری‌های نوین می‌شود. نهادهای نظارتی نیز باید اختیار کافی برای الزامات اصلاحی و اعمال تحریم‌های قانونی داشته باشند. در نهایت، نظارت حقوقی مؤثر،

هوشمند تلقی می‌گردد. پژوهش‌های داخلی بر این نکته تأکید کرده‌اند که عدم افشای کافی اطلاعات درباره سازوکار تصمیم‌گیری می‌تواند موجب بطلان یا قابلیت فسخ قرارداد گردد (فاضلی و همکاران، ۱۴۰۰: ۱۲۱). در پرتو چنین نگاهی، جایگاه کاربر در روابط قراردادی از وضعیت منفعل به موقعیتی فعال و متوازن تغییر خواهد یافت. این تحول با اصول انصاف قراردادی و حمایت از طرف ضعیف در حقوق مدنی همساز است.

۶- جایگاه کاربر در نظام‌های حمایت از داده و نظارت بر تصمیمات الگوریتمی

در نظام‌های نوین حمایت از داده، کاربر به‌عنوان صاحب حق حاکمیت بر داده‌های شخصی شناخته می‌شود و این شناسایی، مبنای تنظیم‌گری در حوزه تصمیمات الگوریتمی را شکل می‌دهد. اصول بنیادینی همچون مشروعیت پردازش، ضرورت رضایت آگاهانه و ممنوعیت جمع‌آوری داده بیش‌ازحد، محور اصلی نظام‌های حقوقی در اروپا، آمریکا و سایر کشورها را تشکیل می‌دهد. مقررات عمومی حفاظت از داده اتحادیه اروپا (GDPR) با وضع قواعد سخت‌گیرانه درباره پردازش خودکار، جایگاه کاربر را به‌صورت تقنینی تقویت کرده است. این مقررات نشان می‌دهند که تصمیمات خودکار بدون شفافیت کافی می‌توانند به تضعیف حقوق کاربران و اختلال در رفاه آنان منجر شوند (Cheong, 2024: 4). نظام‌های ملی نیز بر ضرورت رعایت استانداردهای امنیتی و لزوم ایجاد سازوکارهای مؤثر برای محدودسازی پردازش تأکید دارند. در ایران، هرچند قانون جامع حمایت از داده وجود ندارد، اما اصولی در قوانین تجارت الکترونیک و آیین‌نامه‌های حوزه ارتباطات، حداقل‌سازی جمع‌آوری داده و ضرورت اخذ رضایت معتبر را پیش‌بینی کرده‌اند. پژوهش‌های تطبیقی بیان می‌کنند که عدم شفافیت در مدل‌های خودکار به نقض حقوق بنیادین کاربران و ایجاد تبعیض ساختاری منجر می‌شود (Narayanan & Kapoor, 2024: 56-58).

در این چارچوب، حق دسترسی کاربر به منطق تصمیم، حق اصلاح داده و حق فراموش‌شدن از ارکان اساسی حمایت از داده محسوب می‌شود. علاوه بر این، تضمین امنیت داده و ایجاد لایه‌های فنی برای جلوگیری از دسترسی غیرمجاز از الزامات بنیادین حمایت حقوقی است. هرگونه پردازش بدون مبنای معتبر می‌تواند مسئولیت مدنی و حتی کیفری برای ارائه‌دهندگان خدمات ایجاد کند. نظام‌های حمایتی همچنین تکلیف می‌کنند که کاربران بتوانند تصمیمات خودکار را به چالش بکشند. اجرای

اساس مشروعیت تصمیمات الگوریتمی و تضمین‌کننده حقوق کاربران در جامعه دیجیتال است.

۷- جایگاه فراملی حقوق کاربران در نظام تصمیم‌گیری الگوریتمی

در سطح حقوق بشر بین‌الملل، معاهدات الزام‌آور سازمان ملل چارچوبی روشن برای صیانت از کاربران در برابر تصمیمات خودکار ایجاد کرده‌اند. ماده ۱ اعلامیه جهانی حقوق بشر بر کرامت انسانی تأکید دارد و این اصل باید در بستر تصمیمات الگوریتمی نیز رعایت شود. ماده ۲ و ۲۶ میثاق حقوق مدنی و سیاسی دولت‌ها را ملزم می‌کند که از هرگونه تبعیض ناشی از پردازش خودکار داده‌ها جلوگیری کنند. ماده ۱۴ همان میثاق نیز حق بر دادرسی عادلانه را تضمین می‌کند و این حق شامل اعتراض به تصمیمات الگوریتمی است. همچنین ماده ۶ میثاق حقوق اقتصادی، اجتماعی و فرهنگی حق برخورداری از فرصت برابر در فعالیت‌های اقتصادی را مقرر داشته و این حق در برابر تصمیمات شغلی مبتنی بر هوش مصنوعی نیز قابل استناد است. گزارش‌های کمیسر عالی حقوق بشر توسعه هوش مصنوعی را مقید به رعایت اصول شفافیت و پاسخگویی می‌دانند. یافته‌های تجربی نشان می‌دهد که تصمیمات قضایی مبتنی بر الگوریتم می‌تواند زمینه تبعیض ساختاری را فراهم کند (Herrera-Tapias et al, 2025: 1229).

این اسناد تأکید دارند که دولت‌ها باید امکان بازبینی انسانی مؤثر را تضمین نمایند. همچنین لازم است دسترسی کاربران به منطق تصمیم در قالب توضیح‌پذیری کافی فراهم گردد. دولت‌ها مکلف‌اند مکانیسم‌های جبرانی سریع و مؤثر برای کاربران ایجاد کنند. فقدان این مکانیسم‌ها می‌تواند اصل برابری سلاح‌ها در دادرسی را تضعیف کند. معاهدات سازمان ملل همچنین دولت‌ها را از انتقال ناموجه ریسک به کاربران منع می‌کنند. این معاهدات بر لزوم ارزیابی اثرات حقوق بشری در سامانه‌های پرخطر تأکید دارند. نهایتاً، قواعد سازمان ملل بنیانی الزام‌آور برای تضمین حقوق کاربران در عصر تصمیم‌گیری الگوریتمی فراهم می‌آورند. در سطح فراملی، اسناد و قواعد آن‌سیترا ل نقش مهمی در تنظیم حقوق کاربران در برابر تصمیمات خودکار ایفا می‌کنند. «یادداشت‌های آن‌سیترا ل درباره تجارت الکترونیک و سامانه‌های خودکار» مقرر می‌کند که تصمیمات الگوریتمی نباید موجب محرومیت اشخاص از حقوق قراردادی بدون امکان اعتراض مؤثر شود. این سند همچنین بر لزوم تضمین قابلیت انتساب و وجود

کنترل انسانی در فرآیندهای خودکار تأکید دارد. اصول کلی آن‌سیترا ل در حوزه هوش مصنوعی بر ضرورت شفافیت عملیاتی و ثبت سوابق تصمیمات استوار است. این اصول به دولت‌ها توصیه می‌کنند که نسبت به طراحی سازوکارهای پاسخگویی فنی و حقوقی اقدام کنند. علاوه بر این، قواعد سازمان ملل در قطعنامه‌های شورای حقوق بشر بر لزوم حفاظت مؤثر از حقوق بنیادین در برابر فناوری‌های نوین تأکید دارند. قطعنامه ۲۳/۴۷ شورا تأکید می‌کند که هیچ سامانه هوش مصنوعی نباید موجب نقض حقوق بنیادین افراد شود. یافته‌های اخلاقی نیز نشان می‌دهد که صیانت از کرامت انسانی مستلزم اعمال نظارت انسانی مداوم است (Mura & Stehlíková, 2025: 488).

بر اساس این چارچوب‌ها، دولت‌ها باید ارزیابی اثرات حقوق بشری را پیش از استقرار سامانه‌های پرخطر انجام دهند. همچنین لازم است مسیرهای جبران خسارت برای کاربران مطابق استانداردهای بین‌المللی طراحی گردد. فقدان سازوکارهای شکایت می‌تواند موجب نقض حق بر دسترسی به عدالت شود. اسناد آن‌سیترا ل همچنین بر ضرورت تنظیم‌گری مبتنی بر ریسک تأکید دارند. این تنظیم‌گری باید شدت نظارت را با میزان خطر سامانه متناسب سازد. مسئولیت توسعه‌دهندگان و بهره‌برداران نیز باید بر اساس میزان کنترل قابل اعمال تعیین شود. در نهایت، قواعد سازمان ملل و آن‌سیترا ل شبکه‌ای هماهنگ از تعهدات را برای حمایت از کاربران در برابر تصمیمات الگوریتمی ایجاد می‌کنند.

۸- چالش‌های اثبات و انتساب مسئولیت در الگوریتم‌های یادگیرنده

در نظام سنتی مسئولیت مدنی، زیان‌دیده مکلف است رابطه علیت میان فعل زیان‌بار و ضرر را اثبات نماید، اما در بستر الگوریتم‌های یادگیرنده این امر با دشواری‌های اساسی مواجه است. یادگیری مستمر سامانه موجب تغییر قواعد تصمیم‌گیری می‌گردد و امکان بازسازی مسیر علیت به‌صورت شفاف و خطی را از میان می‌برد. در نتیجه، اثبات تقصیر یا حتی اثبات وقوع خطای فنی با موانع جدی روبه‌رو می‌شود. این وضعیت به‌ویژه در دعاوی مدنی سبب می‌گردد که بار اثبات غیرمتعارفی بر دوش زیان‌دیده تحمیل شود. به همین دلیل، برخی نظام‌های حقوقی پیشنهاد کرده‌اند که بار اثبات از زیان‌دیده به بهره‌بردار یا طراح منتقل گردد. چنین رویکردی نه تنها کارایی جبرانی نهاد مسئولیت را تقویت می‌کند بلکه مانع از بی‌اثر شدن حق دسترسی به عدالت

اثبات دعوی ناکارآمد خواهد شد. از این‌رو، برخی نظام‌ها پیشنهاد کرده‌اند که الزام به ثبت و نگهداری سوابق تصمیمات الگوریتمی به‌عنوان تکلیف قانونی مقرر گردد. چنین الزامی می‌تواند به زیان‌دیده امکان دهد تا شواهد فنی لازم برای اقامه دعوی را به دست آورند. موسوی بر این نکته تصریح می‌کند که بدون تضمین شفافیت، نظام مسئولیت مدنی در حوزه هوش مصنوعی عملاً بی‌اثر خواهد شد (موسوی، ۱۳۹۹: ۵۹). در پرتو چنین دیدگاهی، الزام به شفافیت فنی بخشی از تکالیف عمومی ارائه‌دهندگان سامانه محسوب می‌شود. بدین ترتیب، رفع مانع شفافیت یکی از پیش‌شرط‌های تحقق اثبات و انتساب مسئولیت در الگوریتم‌های یادگیرنده است.

یکی از چالش‌های مهم در انتساب مسئولیت به الگوریتم‌های یادگیرنده، مسأله پیش‌بینی‌ناپذیری رفتار سامانه است که ناشی از ماهیت پویا و خودسازگار آن می‌باشد. در این چارچوب، دادگاه‌ها برای احراز تقصیر یا رابطه علیت ناگزیرند به استانداردهای فنی و حرفه‌ای متوسل شوند، درحالی‌که این استانداردها همواره در حال تحول هستند. این وضعیت به بی‌ثباتی هنجاری منجر می‌شود و پیش‌بینی‌پذیری حقوقی را کاهش می‌دهد. از منظر عدالت قضایی، چنین وضعیتی می‌تواند به نفع توسعه‌دهندگان و به زیان زیان‌دیده عمل کند، زیرا آنان قادر به اثبات نقض تعهدات فنی نخواهند بود. بنابراین، پیشنهاد شده است که معیار «رفتار متعارف توسعه‌دهنده هوش مصنوعی» به‌عنوان مبنای سنجش تقصیر پذیرفته شود. این معیار به‌جای تمرکز بر پیش‌بینی دقیق خروجی، بر رعایت موازین احتیاطی و طراحی ایمن تأکید می‌نماید. در پژوهش‌های تطبیقی بر این نکته تأکید شده است که توسعه استانداردهای مسئولیت‌محور، شرط لازم برای امکان انتساب مسئولیت در الگوریتم‌های یادگیرنده است (Calo, 2015: 189). چنین معیاری ضمن کاهش ابهام در رویه قضایی، امکان انطباق نهاد مسئولیت با تحولات فناورانه را فراهم می‌سازد. از این‌رو، پذیرش معیار رفتاری جدید گامی اساسی در جهت کارآمدی نظام مسئولیت مدنی خواهد بود.

بعد دیگر چالش‌های انتساب مسئولیت، ماهیت جمعی و غیرشخصی داده‌هایی است که الگوریتم‌های یادگیرنده بر پایه‌ی آن‌ها عمل می‌کنند و این ویژگی موجب می‌شود که نسبت‌دادن خروجی الگوریتم به یک شخص معین از لحاظ حقوقی با دشواری‌های اساسی مواجه گردد. الگوریتم‌هایی که بر داده‌های اشتراکی و چندمنبعی تکیه دارند، اسبابی را ایجاد می‌کنند که

می‌شود. در پژوهش‌های داخلی نیز تأکید شده است که جابه‌جایی بار اثبات در حوزه الگوریتم‌های خودآموز ضرورتی اجتناب‌ناپذیر برای تضمین حمایت مؤثر از زیان‌دیده است (صالحی، ۱۴۰۰: ۷۶). این تغییر رویکرد در واقع پاسخی به پیچیدگی‌های فنی و عدم تقارن اطلاعاتی میان کاربر و ارائه‌دهنده است. بر این اساس، عدالت قضایی اقتضا می‌کند که بار اثبات به شکلی عادلانه باز توزیع گردد. بدین ترتیب، اثبات رابطه علیت در الگوریتم‌های یادگیرنده نیازمند اصلاح قواعد سنتی و تطبیق با مقتضیات فناورانه است.

یکی از دیگر چالش‌های بنیادین در زمینه انتساب مسئولیت، تعیین شخص یا اشخاصی است که باید در برابر زیان‌دیده پاسخگو باشند. الگوریتم‌های یادگیرنده نه‌تنها بر مبنای طراحی اولیه بلکه بر اساس داده‌های ورودی و تعاملات مستمر تکامل می‌یابند و این امر باعث می‌شود سهم دقیق هر یک از بازیگران در وقوع زیان مبهم گردد. در چنین شرایطی، استفاده از معیارهای سنتی تقصیر یا مباشرت در فعل زیان‌بار کارآمدی خود را از دست می‌دهد. افزون بر این، دخالت چندین حلقه از زنجیره تولید، توسعه و بهره‌برداری امکان انتساب خطا به یک عامل منفرد را دشوار می‌سازد (Christian, 2020: 54-55). به همین دلیل، نظریه‌های نوین بر ضرورت اتخاذ رویکرد نهادی و شبکه‌ای در تعیین مسئولیت تأکید می‌کنند. در این رویکرد، هر حلقه بر اساس نقش و میزان کنترل خود در معرض مسئولیت قرار می‌گیرد. پژوهش‌های داخلی اخیر نشان داده‌اند که چنین تقسیم مسئولیتی می‌تواند از گریز اشخاص از بار پاسخگویی جلوگیری نماید (آذری و همکاران، ۱۴۰۱: ۱۳۳). این الگو در عین حال باید به‌گونه‌ای طراحی گردد که امکان جبران سریع و مؤثر خسارت برای بزه‌دیده فراهم شود. بنابراین، بازتعریف مفهوم انتساب مسئولیت در الگوریتم‌های یادگیرنده شرط اساسی کارآمدی نهاد مسئولیت مدنی محسوب می‌شود.

چالش مهم دیگر، مسأله شفافیت و قابلیت دسترسی به اطلاعات مربوط به فرایند تصمیم‌گیری الگوریتمی است که به‌طور مستقیم بر امکان اثبات دعوی مدنی تأثیر می‌گذارد. الگوریتم‌های یادگیرنده اغلب با ساختارهای موسوم به «جعبه سیاه» توصیف می‌شوند که در آن نه کاربران و نه حتی توسعه‌دهندگان توانایی درک کامل منطق تصمیم را ندارند. این امر سبب می‌شود زیان‌دیده نتواند دلایل کافی برای انتساب تقصیر یا بی‌مبالاتی ارائه‌دهنده گردآوری نماید. در غیاب شفافیت، قواعد سنتی ادله

منشأ واحد آن قابل احراز نیست و بنابراین مسئولیت باید در قالب ساختارهای نهادی بازتعریف شود. چنین بازتعریفی به دولت‌ها تکلیف می‌کند که چارچوب‌های قانونی و جبرانی ویژه برای حمایت از زیان‌دیدگان از تصمیمات خودکار ایجاد کنند. افزون بر این، ایجاد صندوق‌های جبرانی و مکانیسم‌های بیمه‌ای مکمل برای توزیع منصفانه ریسک میان بازیگران مختلف زنجیره فناوری ضروری دانسته شده است. این راهکارها ضمن تأمین منافع زیان‌دیدگان، مانع از گریز توسعه‌دهندگان و بهره‌برداران از بار مسئولیت می‌گردد. پژوهش‌های تطبیقی نشان داده‌اند که باز توزیع نهادی مسئولیت و شناسایی مکانیسم‌های جبرانی جمعی، کارآمدترین شیوه مواجهه با خطرات نوظهور فناوری‌های یادگیرنده است (Erdélyi & Goldsmith, 2021: 152–153).

اهمیت این الگو در حوزه‌های پرخطر مانند مراقبت‌های سلامت و حمل‌ونقل هوشمند دو چندان است، زیرا میزان خسارات بالقوه می‌تواند فراتر از ظرفیت جبران فردی باشد. با این حال، در وضعیتی که بیمه و صندوق‌های جبرانی وجود نداشته یا از پوشش کافی برخوردار نباشند، نظام حقوقی باید به‌سوی پذیرش «مسئولیت قانونی مستقیم دولت» یا «مسئولیت تضامنی اجباری توسعه‌دهندگان» حرکت کند. چنین سازوکاری می‌تواند از طریق وضع قوانین حمایتی و الزام‌آور تضمین شود تا خلأ نهادی جبران گردد. علاوه بر این، امکان ایجاد «تعهد مراقبت تقویت‌شده» برای بهره‌برداران و طراحان نیز مطرح است تا بار احتیاط فنی افزایش یابد. در موارد خاص، اعمال قواعد ویژه مسئولیت بدون تقصیر نیز می‌تواند ضرورت یابد تا سنگینی بار اثبات از دوش زیان‌دیده برداشته شود. این رویکرد هم‌زمان با اصول عدالت اجتماعی و کارایی اقتصادی هماهنگ است و از تضعیف اعتماد عمومی به سامانه‌های هوش مصنوعی جلوگیری می‌کند. در نتیجه، انتساب مسئولیت در این حوزه باید در قالب الگوی چندلایه و نهادی صورت گیرد تا خلأهای بیمه‌ای یا نهادی مانع جبران خسارت نگردد.

۹- راهکارهای حقوقی و تنظیم‌گری در حوزه تصمیمات هوش مصنوعی

نخستین راهکار بنیادین در تنظیم‌گری حوزه تصمیمات هوش مصنوعی، ایجاد چارچوب‌های الزام‌آور حقوقی است که بتوانند ماهیت پویای این فناوری را پوشش دهند. این چارچوب‌ها باید بر اصولی چون شفافیت، پاسخگویی و قابلیت پیش‌بینی استوار باشند تا از خلأهای هنجاری در فرآیند تصمیم‌گیری جلوگیری

گردد. در این مسیر، تکیه صرف بر توصیه‌های اخلاقی کافی نیست و باید ضمانت اجرایی حقوقی لازم برای تخلف از استانداردها پیش‌بینی شود. در حقوق تطبیقی تأکید می‌شود که قواعد الزام‌آور حقوقی در کنار ابزارهای نظارتی می‌توانند موجب تضمین رعایت حقوق بنیادین کاربران گردند (Surden, 2019: 1310). یکی از ابزارهای اساسی در این چارچوب، ایجاد نهادهای تخصصی نظارتی است که صلاحیت فنی و حقوقی لازم برای ممیزی الگوریتم‌ها را دارا باشند. چنین نهادهایی می‌توانند استقلال لازم برای کنترل فرآیندهای تصمیم‌گیری و تضمین بی‌طرفی در ارزیابی عملکرد سامانه‌ها را فراهم کنند. علاوه بر آن، الزامات مربوط به ثبت و مستندسازی تصمیمات الگوریتمی باید در سطح تقنینی پیش‌بینی شوند تا امکان بازبینی قضایی فراهم آید. با این حال، تنظیم‌گری نباید مانع نوآوری مشروع گردد بلکه باید تعادلی میان حمایت از حقوق بنیادین و توسعه اقتصادی برقرار نماید. پژوهش‌های اخیر نشان داده‌اند که طراحی قواعد منعطف و مبتنی بر ریسک بهترین الگو برای تنظیم‌گری در این حوزه محسوب می‌شود (Edwards, 2021: 198). در نهایت، این رویکرد می‌تواند به تثبیت مشروعیت حقوقی تصمیمات الگوریتمی و افزایش اعتماد عمومی به کارکردهای هوش مصنوعی منجر شود.

راهکار دیگر، تلفیق قواعد سنتی حقوق خصوصی با الزامات خاص فناوری‌های نوین است که به ایجاد نظامی ترکیبی و کارآمد منجر می‌گردد. قواعدی مانند حسن نیت، تعهد به اطلاع‌رسانی و مسئولیت مدنی باید در پرتو شرایط جدید بازتفسیر شوند. این بازتفسیر باید به گونه‌ای باشد که بتواند تعهدات قراردادی و غیر قراردادی ارائه‌دهندگان خدمات را به‌طور هم‌زمان پوشش دهد. در واقع، کارآمدی نهاد مسئولیت در گرو تطبیق با ماهیت خاص تصمیمات خودکار و خطرات ناشی از آن است. در این چارچوب، بهره‌گیری از ظرفیت‌های فقهی و حقوق داخلی می‌تواند نقش مهمی در بومی‌سازی قواعد ایفا نماید. افزون بر آن، باید راهکارهایی برای تضمین حق اعتراض و جبران خسارت کاربران پیش‌بینی گردد تا عدالت ترمیمی به‌طور واقعی تحقق یابد. چنین تدابیری مانع از آن می‌شوند که کاربران در موقعیتی نابرابر در برابر ارائه‌دهندگان خدمات قرار گیرند. باز تعریف این اصول در حوزه الگوریتم‌های هوشمند موجب تقویت اعتماد اجتماعی و کاهش خطرات ناشی از سوءاستفاده از فناوری خواهد شد. برخی پژوهش‌های داخلی تصریح کرده‌اند که بدون بازاندیشی در اصول

داشته باشند. چنین نهادهایی می‌توانند با همکاری دولت‌ها و بخش خصوصی، مانع از بروز شکاف‌های نظارتی شوند. از منظر حقوق بین‌الملل، تنظیم‌گری هماهنگ علاوه بر تضمین حمایت از حقوق بشر دیجیتال، موجب تسهیل تجارت و نوآوری جهانی نیز خواهد شد. این هماهنگی به‌ویژه در حوزه‌هایی مانند تجارت الکترونیک و خدمات برخط اهمیتی مضاعف دارد. پژوهش‌های تطبیقی تصریح کرده‌اند که همکاری‌های بین‌المللی در زمینه تنظیم‌گری هوش مصنوعی شرط لازم برای انسجام و مشروعیت این نظام حقوقی در آینده است (Vujicic, 2024: 3460-3465). بر این اساس، تنها با ایجاد قواعد هماهنگ فراملی است که می‌توان از تعارضات حقوقی و تضییع حقوق کاربران در فضای جهانی جلوگیری کرد.

نتیجه‌گیری

این پژوهش نشان داد که چارچوب‌های سنتی مسئولیت مدنی در برابر تصمیمات الگوریتمی توان پاسخگویی کافی ندارند. معیارهای کلاسیک تقصیر، رابطه سببیت و پیش‌بینی‌پذیری دیگر قادر به تبیین پویایی سامانه‌های خودآموز نیستند. بررسی‌ها آشکار ساخت که ساختار تصمیم‌گیری الگوریتمی غالباً غیر شفاف و پیچیده است. همین امر مسیر تحلیل مسئولیت را دشوارتر از گذشته می‌سازد. یافته‌ها نشان داد که حتی طراحان و توسعه‌دهندگان نیز به منطق درونی تصمیمات سامانه‌ها احاطه کامل ندارند. این ناتوانی در فهم کامل فرآیند تصمیم‌گیری، انتساب تقصیر مستقیم را با ابهام مواجه می‌کند. علاوه بر این، عدم توضیح‌پذیری الگوریتم‌ها مانع از ارائه ادله کافی توسط زیان‌دیدگان می‌شود. تداوم چنین وضعیتی می‌تواند حق جبران خسارت را عملاً محدود سازد. داده‌های پژوهش ثابت کرد که قواعد سنتی حقوقی بدون بازآفرینی قادر به مدیریت ریسک‌های نوظهور نیستند. بر این اساس، ضرورت اصلاح مفاهیم بنیادین مسئولیت بیش‌ازپیش آشکار می‌شود.

تحلیل تطبیقی نشان داد که نظام‌های حقوقی پیشرو به سمت پذیرش الگوهای چندعاملی حرکت کرده‌اند. این الگوها امکان تقسیم عادلانه بار زیان میان چندین بازیگر را فراهم می‌کنند. بررسی مقررات اتحادیه اروپا نیز بر اهمیت ایجاد قواعد الزام‌آور مستقل برای فناوری‌های پر ریسک تأکید دارد. این تجربه‌ها بیانگر آن است که مسئولیت باید از سطح فردی به سطح نهادی گسترش یابد. پژوهش حاضر نشان داد که مسئولیت شبکه‌ای می‌تواند سازوکار مؤثرتری برای مدیریت ریسک باشد. این مدل

سنتی، تنظیم‌گری در این حوزه عملاً ناقص خواهد بود (بحر کاظمی و محمودی، ۱۴۰۰: ۶۰). بدین ترتیب، ترکیب قواعد سنتی و نوین ابزار مناسبی برای برقراری توازن میان حقوق کاربران و منافع اقتصادی خواهد بود.

یکی دیگر از راهکارهای اساسی در تنظیم‌گری تصمیمات هوش مصنوعی، پذیرش الگوی «تنظیم‌گری مبتنی بر ریسک» است که طی آن شدت مداخله حقوقی بر مبنای سطح خطر سامانه تعیین می‌گردد. بدین معنا که سامانه‌های پرخطر مانند حوزه سلامت، حمل‌ونقل یا عدالت کیفری نیازمند قواعد الزام‌آور سخت‌گیرانه خواهند بود، درحالی‌که سامانه‌های کم‌خطر می‌توانند از الزامات انعطاف‌پذیرتری برخوردار شوند. این رویکرد علاوه بر کارآمدی، موجب تخصیص بهینه منابع نظارتی نیز می‌گردد. در چنین الگویی، تعادل میان حمایت از حقوق بنیادین شهروندان و توسعه نوآوری به بهترین شکل محقق می‌شود. افزون بر آن، این الگو قابلیت انطباق با تحولات فناورانه را دارد، زیرا می‌تواند در پرتو ارزیابی‌های مستمر بازبینی گردد. همچنین، این نظام به دولت‌ها امکان می‌دهد تا با طبقه‌بندی فناوری‌ها، ابزارهای حقوقی متناسب با هر سطح خطر را به کار گیرند. از حیث حقوقی، چنین رویکردی مانع از یکسان‌نگاری غیرمنصفانه میان فناوری‌های ناهمگون خواهد شد. پژوهش‌های تطبیقی نشان داده است که تنظیم‌گری مبتنی بر ریسک به‌عنوان مؤثرترین راهبرد برای مواجهه با پیچیدگی‌های هوش مصنوعی شناخته شده است (Veale & Borgesius, 2021: 280). این امر نشان می‌دهد که انعطاف حقوقی در عین سخت‌گیری گزینشی، شرط لازم برای تضمین مشروعیت و اثربخشی نظام تنظیم‌گری است. بدین ترتیب، الگوریتم‌ها می‌توانند هم در مسیر نوآوری و هم در چهارچوب حقوق بنیادین هدایت شوند.

راهکار مکمل دیگر، طراحی سازوکارهای بین‌المللی هماهنگ برای تنظیم‌گری تصمیمات الگوریتمی است تا از تعارضات حقوقی فراملی جلوگیری گردد. هوش مصنوعی ذاتاً پدیده‌ای فرامرزی است و محدود کردن آن به قواعد داخلی نمی‌تواند پاسخگوی چالش‌های جهانی آن باشد. در این زمینه، ضرورت دارد اصول عامی در سطح بین‌الملل تدوین گردد که به‌مثابه «حداقل استانداردهای جهانی» عمل کنند. این اصول می‌توانند ناظر بر شفافیت، منع تبعیض، تضمین حریم خصوصی و حق دسترسی به جبران خسارت باشند. علاوه بر این، نهادهای بین‌المللی باید صلاحیت نظارت و ارزیابی بر اجرای این اصول را

ظرفیت بیشتری برای حمایت از بزه‌دیدگان ایجاد می‌کند. شواهد ارائه‌شده در مقاله این روند تحول را تأیید کرد. یافته‌ها نشان داد که فناوری‌های نوین نیازمند استانداردهای شفافیت بالاتری هستند. این الزام بخشی از ساختار جدید مسئولیت محسوب می‌شود. در نتیجه، فرضیه اصلی مقاله مبنی بر ناکارآمدی چارچوب‌های سنتی به‌طور کامل اثبات گردید.

این مقاله همچنین جایگاه کاربران در فرآیند تصمیم‌گیری الگوریتمی را واکاوی کرد و نشان داد که کاربران در معرض تأثیر تصمیمات خودکار قرار دارند بدون آنکه همواره از سازوکار درونی این تصمیمات آگاهی داشته باشند. این وضعیت موجب شده است که رضایت آگاهانه، که یکی از ارکان اساسی روابط حقوقی است، در بسیاری از موارد تحقق نیابد. یافته‌ها ثابت کرد که افشای ناکافی اطلاعات و پیچیدگی فنی سامانه‌ها موجب نابرابری قراردادی میان کاربر و ارائه‌دهنده می‌شود. بررسی منابع تطبیقی نشان داد که حقوق بنیادینی مانند حق دسترسی به منطق تصمیم، حق اعتراض و حق بازبینی انسانی باید به‌عنوان مؤلفه‌های ضروری حمایت از کاربران شناخته شوند. این حقوق بخشی از رویکرد «حمایت‌محور» هستند؛ رویکردی که هدف آن ارتقاء جایگاه کاربران از وضعیت مصرف‌کننده منفعل به صاحب حق فعال است. این رویکرد حکم می‌کند که کاربر نه صرفاً بهره‌بردار خدمات بلکه بخشی از چرخه تصمیم‌گیری تلقی شود.

تحلیل انجام‌شده نشان داد که شفافیت الگوریتمی شرط لازم برای تحقق این رویکرد حمایتی است. تجربه اتحادیه اروپا نشان داد که تنظیم‌گری مؤثر زمانی شکل می‌گیرد که حقوق کاربران به‌صورت قواعد الزام‌آور تدوین شود. این یافته بیان می‌کند که حمایت قانونی باید فراتر از توصیه‌های اخلاقی باشد. پژوهش نشان داد که فقدان شفافیت، کاربر را در موضع ضعف ساختاری قرار می‌دهد. این ضعف می‌تواند زمینه‌ساز تبعیض و نقض عدالت شود. بنابراین، تقویت حقوق کاربران جزئی جدایی‌ناپذیر از یک نظام حقوقی کارآمد در عصر هوش مصنوعی است. تحلیل داده‌های پژوهش فرضیه ضرورت ارتقاء جایگاه کاربران به سطح صاحبان حقوق بنیادین را تأیید کرد. این نتیجه نشان می‌دهد که عدالت دیجیتال بدون شناسایی این حقوق قابل تحقق نیست. در نهایت، یافته‌ها آشکار ساخت که پذیرش رویکرد حمایتی برای جلوگیری از تمرکز قدرت در دست سامانه‌های هوشمند و بهره‌برداران آن‌ها ضروری است.

در نهایت، این پژوهش نشان داد که مواجهه با چالش‌های حقوقی هوش مصنوعی مستلزم اتخاذ «رویکردی چندلایه» است؛ رویکردی که بر تعامل و هم‌افزایی میان چند نظام حقوقی، نهادی و حمایتی استوار است. در این تعریف، رویکرد چندلایه به معنای ترکیب هم‌زمان قواعد مسئولیت مدنی، مقررات تنظیم‌گری، سازوکارهای حمایتی و همکاری‌های فراملی است. این رویکرد بر آن است که هیچ‌یک از ابزارهای حقوقی موجود به‌تنهایی قادر به مهار ریسک‌های نوظهور نیستند. نخستین لایه این رویکرد، بازآفرینی قواعد داخلی مسئولیت مدنی است به‌گونه‌ای که بتوانند با استانداردهای نوین شفافیت و پاسخگویی سازگار شوند. دومین لایه، طراحی ابزارهای حمایتی مانند بیمه اجباری و صندوق‌های جبرانی است که بار مالی زیان‌ها را توزیع می‌کنند. سومین لایه، تنظیم‌گری مبتنی بر ریسک است که شدت مداخله دولت را بر اساس سطح خطر فناوری تعیین می‌کند. چهارمین لایه، همکاری بین‌المللی در قالب استانداردهای هماهنگ جهانی است که مانع فرار تنظیمی و تعارض قوانین می‌شود. پنجمین لایه، تقسیم نهادی مسئولیت است که بار جبران را میان طراحان، توسعه‌دهندگان و بهره‌برداران توزیع می‌کند. ششمین لایه، بازتعریف اصول حقوق کیفری با تمرکز بر قصور ساختاری است. این لایه نشان می‌دهد که مسئولیت کیفری در عصر هوش مصنوعی بر ترک مراقبت نهادی تکیه دارد. در مجموع، رویکرد چندلایه یک چارچوب جامع برای مواجهه با ریسک‌های هوش مصنوعی ارائه می‌دهد. این رویکرد از یک‌سو مانع از تمرکز بار مسئولیت بر دوش یک بازیگر می‌شود و از سوی دیگر، عدالت جبرانی را تقویت می‌کند. یافته‌های پژوهش ثابت کرد که تنها در پرتو چنین رویکردی می‌توان نظامی پایدار، متوازن و کارآمد برای تنظیم‌گری ایجاد کرد. این نتیجه بیان می‌کند که آینده مسئولیت حقوقی در حوزه هوش مصنوعی وابسته به پذیرش این چارچوب چندبعدی است. در نهایت، پژوهش نشان داد که رویکرد چندلایه نه‌تنها قابلیت حل اختلافات موجود را دارد، بلکه توانایی پیش‌بینی و مدیریت ریسک‌های آتی را نیز فراهم می‌سازد.

ملاحظات اخلاقی: در تمام مراحل نگارش پژوهش حاضر، صداقت و امانت‌داری رعایت شده است.

تعارض منافع: در این مقاله هیچ‌گونه تضاد منافی وجود ندارد.

سهم نویسندگان: نگارش مقاله مشترکاً توسط نویسندگان صورت گرفته است.

- علی‌پناهی، مریم؛ نصیران نجف‌آبادی، داوود؛ و شیرانی، مسعود (۱۴۰۳). «مسئولیت مدنی ناشی از استفاده هوش مصنوعی در اتحادیه اروپا». *فصلنامه علمی مطالعات فقه اقتصادی*، ۶ (۵): ۱۸-۱.

- علی‌نقیان، اشکان؛ صفدری رنجبر، مصطفی و محمدی، مهدی (۱۴۰۲). «طراحی بسته سیاستی برای توسعه هوش مصنوعی ایران». *سیاست‌گذاری عمومی*، مقالات آماده انتشار.

- کریمی، ماشالله و علی‌آبادی، عبدالصمد (۱۴۰۳). «امکان‌سنجی بهره‌گیری از دادرسی هوش مصنوعی؛ مطالعه فقهی و حقوق ایران». *مطالعات فقه و اصول*، ۷ (۱): ۱۲۷-۱۵۰.

- موسوی، علی؛ حسینی، مهدی و رضایی، فاطمه (۱۴۰۲). «تحولات صنعت رسانه مبتنی بر فناوری هوش مصنوعی: کاربردها، فرصت‌ها و مخاطرات». *فصلنامه مطالعات رسانه‌های نوین*، ۸ (۲): ۱-۲۰.

- ولی‌پور، علی و اسماعیلی، محسن (۱۴۰۰). «امکان‌سنجی مسئولیت مدنی هوش مصنوعی عمومی ناشی از ایجاد ضرر در حقوق مدنی». *اندیشه حقوقی معاصر*، ۲ (۳): ۱-۸.

- یزدانیان، علیرضا (۱۴۰۳). *قواعد عمومی مسئولیت مدنی (جلد اول): با مطالعه تطبیقی در حقوق فرانسه*. تهران: نشر میزان.

ب. منابع انگلیسی

- Birch, J (2024). *The Edge of Sentience: Risk and Precaution in Humans, Other Animals, and AI*. Oxford University Press.

<https://doi.org/10.1093/9780191966729.001.0001>

- Burri, T (2023). A Challenge for Law and Artificial Intelligence. *Nature Machine Intelligence*. Advance Online Publication.

- Calo, R (2016). Robotics and the Lessons of Cyberlaw. In P. Ohm & M. Steenson (Eds.), *Robot Law* (pp. 35-52). Edward Elgar Publishing. <https://doi.org/10.4337/9781783476732>

Cheong, B. C (2024). "Transparency and Accountability in AI Systems: Safeguarding Wellbeing in the Age of Algorithmic Decision-Making". *Frontiers in Human Dynamics*, 6, Article 1421273.

<https://doi.org/10.3389/fhumd.2024.1421273>

- Chesterman, S (2020). "Artificial Intelligence and the Limits of Legal Personality". *International &*

تشکر و قدردانی: از کلیه کسانی که در معرفی منابع و تهیه این مقاله ما را یاری رساندند، تشکر و قدردانی می‌نمایم.

تأمین اعتبار پژوهش: این پژوهش فاقد تأمین‌کننده مالی بوده است.

منابع

الف. منابع فارسی

- اکبری، محمدرضا؛ عباسی، امیر و علوی، احمد (۱۴۰۲). «مسئولیت مدنی ناشی از کاربرد سامانه‌های هوشمند در حقوق ایران». *فصلنامه مطالعات نوین حقوق*، ۹ (۳): ۸۳-۹۴.

- آقا امینی فشمی، حامد و حسینی‌مقدم، سید حسن (۱۴۰۲). «بررسی تطبیقی چالش‌های حقوقی هوش مصنوعی در مراقبت‌های بهداشتی در حقوق ایران و اتحادیه اروپا». *فصلنامه مطالعات تطبیقی حقوق اسلام و غرب*، ۱۱ (۲): ۴۷-۶۸.

- تشخیص، زهرا (۱۴۰۰). «مقدمه‌ای بر چالش‌های هوش مصنوعی در حوزه مسئولیت مدنی». *مجله علمی حقوق خصوصی*، ۱۸ (۱): ۲۵۰-۲۲۷.

- حاجی‌اسماعیلی، میلاد (۱۴۰۳). «چالش‌های مسئولیت مدنی هوش مصنوعی در نظام حقوقی ایران؛ با نگاهی به مقررات‌گذاری در اتحادیه اروپا». *نشریه علمی دولت و حقوق*، ۵ (۱): ۸۱-۹۸.

- حکمت‌نیا، محمود؛ محمدی، مرتضی و واثقی، محسن (۱۳۹۸). «مسئولیت مدنی ناشی از تولید ربات‌های مبتنی بر هوش مصنوعی خودمختار». *حقوق اسلامی*، ۱۶ (۶۰): ۲۵۸-۲۳۱.

- ذاکری‌نیا، حانیه (۱۴۰۲). «ماهیت و مبنای مسئولیت مدنی ناشی از هوش مصنوعی در حقوق ایران و کشورهای اتحادیه اروپا». *مجله حقوق خصوصی*، ۲۰ (۱): ۱۵۲-۱۳۵.

- رجبی، عبدالله (۱۳۹۸). «ضمان در هوش مصنوعی». *مطالعات حقوق تطبیقی*، ۱۰ (۲): ۴۴۹-۴۶۶.

- سیفی، آناهیتا و رزمخواه، نجمه (۱۴۰۰). «تأمین و اجرایی شدن حق اشتغال زنان در پرتو توسعه و گسترش هوش مصنوعی». *مطالعات حقوق بشر اسلامی*، ۱۰ (۲): ۱۵۳-۱۸۰.

- صابری، مصطفی (۱۴۰۳). «نقش مصالح مرسله در توسعه مسئولیت مدنی در حوزه فناوری‌های نوین». *مجله فقه کاربردی*، ۴ (۱): ۹۰-۱۰۵.

in the Age of AI. Equilibrium”. *Quarterly Journal of Economics and Economic Policy*, 20(2): 479–507. <https://doi.org/10.24136/eq.3743>

- Narayanan, A & Kapoor, S (2024). *AI Snake Oil: What Artificial Intelligence Can Do, What It Can't, and How to Tell the Difference?* Princeton University Press.

- Shneiderman, B (2022). “Responsible AI: Bridging from Ethics to Practice”. *Communications of the ACM*, 65(6): 30–35. <https://doi.org/10.1145/3531000>

- Solum, L. B (2020). *Legal Personhood for Artificial Intelligence*. Oxford University Press.

- Stern, A. D (2022). “AI insurance: How liability Insurance can drive the Adoption of AI”. *NEJM Catalyst Innovations in Care Delivery*, 3(4), Article CAT.21.0242 .

- Surden, H (2019). “Artificial Intelligence and Law: An Overview”. *Georgia State University Law Review*, 35(4): 1305–1370.

- Tai, E. T. T (2022). Liability for AI Decision-Making. In L. A. DiMatteo, C. Poncibò, & M. Cannarsa (Eds.), *The Cambridge Handbook of Artificial Intelligence: Global Perspectives on Law and Ethics* (pp. 116–131). Cambridge University Press. <https://doi.org/10.1017/9781108955680>

- Veale, M & Zuiderveen Borgesius, F (2021). “Demystifying the Draft EU Artificial Intelligence Act—Analysing the Good, the Bad, and the Unclear Elements of the Proposed Approach”. *Computer Law Review International*, 22(4): 97–112.

Vujicic, J (2024). “Global AI Regulation and its Impact on Technology Business: A Comparative Legal Framework Analysis”. *World Journal of Advanced Research and Reviews*, 24(3): 3457–3463.

- Wagner, J & Wagner, H (2020). “Algorithmic Discrimination and Accountability in AI Systems”. *AI & Society*, 35(1): 1–17.

Comparative Law Quarterly, 69(4): 819–844. <https://doi.org/10.1017/S0020589320000366>

- Christian, B (2020). *The Alignment Problem: Machine Learning and Human Values*. W. W. Norton & Company.

- Edwards, L & Veale, M (2017). “Slave to the Algorithm? Why a Right to an Explanation is Probably not the Remedy you are Looking for”. *Duke Law & Technology Review*, 16(1): 18–84.

<https://doi.org/10.2139/ssrn.2972855>

- Erdélyi, O. J & Goldsmith, J (2021). “The AI Liability Puzzle and a Fund-Based Work-Around”. *Journal of Artificial Intelligence Research*, 71: 143–168. <https://doi.org/10.1613/JAIR.1.12551>

- Garon, J. M (2024). “Ethics 3.0 – Attorney Responsibility in the age of Generative AI”. *The Business Lawyer*, 79: 209-220.

- Herrera-Tapias, B. A; Hernández Guzmán, D; Zambam, N. J; Turatti, L; Rodríguez, F. A; Fröhlich, S; Staffen, M. R; Carvajal Muñoz, P; Gutierrez Reyes, A; Soto de la Torre, G; Reyes Duarte, N & Palencia Ramos, E (2025). “Algorithmic Discrimination and Explainable Artificial Intelligence in the Judiciary: A Case Study of the Constitutional Court of Colombia”. *Procedia Computer Science*, 257: 1227–1232.

- Judge, B; Nitzberg, M & Russell, S (2025). “When Code isn't law: Rethinking Regulation for Artificial Intelligence”. *Policy and Society*, 44(1): 85–97. <https://doi.org/10.1093/polsoc/puae020>

- McQuillan, D (2022). *Resisting AI: An Anti-fascist Approach to Artificial Intelligence*. Bristol University Press.

- Mökander, J; Axente, M; Casolari, F & Floridi, L (2022). “Conformity Assessments and Post-Market Monitoring: A Guide to the Role of Auditing in the Proposed European AI Regulation”. *Minds and Machines*, 32(2): 241–268. <https://doi.org/10.1007/s11023-021-09577-4>

- Mura, L & Stehlíková, B (2025). “The Ethics of Artificial Intelligence: Safeguarding Human Dignity, Social Justice and Environmental Stability