

۱-۱۵	صفحات	۲۲	شماره انتشار	۱۴۰۴	سال انتشار
------	-------	----	--------------	------	------------

مطالعه تطبیقی مسئولیت مدنی و سیاست جنایی در قبال خسارات ناشی از خودمختاری هوش مصنوعی؛ با تأکید بر حقوق ایران، اتحادیه اروپا و ایالات متحده

آقای دکتر مهدی شهرکی و خانم نهال ابراهیم پور

دکتری حقوق جزا و جرم شناسی، واحد بین الملل امارات، دانشگاه آزاد اسلامی تهران، ایران.

دکتری گروه حقوق خصوصی، واحد کرج، دانشگاه آزاد اسلامی، کرج، ایران

چکیده

هوش مصنوعی خودمختار با برخورداری از قابلیت یادگیری، تصمیم‌گیری مستقل و تغییر رفتار بر پایه داده‌های جدید، یکی از مهم‌ترین چالش‌های نظام‌های حقوقی معاصر را در حوزه مسئولیت مدنی و سیاست جنایی ایجاد کرده است. در چنین شرایطی، انتساب خسارات ناشی از تصمیمات مستقل سامانه‌های هوش مصنوعی به اشخاص انسانی، اعم از تولیدکننده، توسعه‌دهنده، بهره‌بردار یا مالک، با دشواری‌های جدی مواجه شده و مبانی سنتی مسئولیت مدنی مبتنی بر تقصیر، رابطه سببیت و قابلیت پیش‌بینی زیان، کارآمدی پیشین خود را تا حد زیادی از دست داده‌اند. از سوی دیگر، توسعه کاربردهای هوش مصنوعی در حوزه‌هایی مانند حمل‌ونقل، پزشکی، خدمات مالی، امنیت و قضا، سیاست جنایی دولت‌ها را نیز با ضرورت طراحی سازوکارهای نوین پیشگیرانه، نظارتی و جبرانی روبه‌رو ساخته است. در همین راستا، پژوهش حاضر با هدف تبیین مبانی مسئولیت مدنی ناشی از خودمختاری هوش مصنوعی و تحلیل سیاست جنایی نظام‌های حقوقی ایران، اتحادیه اروپا و ایالات متحده، به روش توصیفی - تحلیلی و با رویکرد تطبیقی انجام شده است. داده‌های پژوهش از طریق مطالعه منابع کتابخانه‌ای، قوانین، اسناد بین‌المللی و مقالات علمی معتبر فارسی و انگلیسی گردآوری و تحلیل شده‌اند. یافته‌های پژوهش نشان می‌دهد که اتحادیه اروپا با تصویب مقررات اختصاصی در حوزه هوش مصنوعی و توسعه نظام مسئولیت مبتنی بر خطر، مسئولیت بدون تقصیر، بیمه اجباری و تسهیل اثبات دعوا، درصدد جبران خلأهای موجود برآمده است؛ در مقابل، نظام حقوقی ایالات متحده همچنان عمدتاً بر قواعد مسئولیت محصول، تقصیر حرفه‌ای و رویه قضایی تکیه دارد، هرچند گرایش به تنظیم‌گری تخصصی نیز در حال افزایش است. در حقوق ایران نیز با وجود امکان استناد به قواعد عام مسئولیت مدنی، از جمله نظریه اتلاف، تسبیب، ضمان ید و قواعد فقهی، فقدان مقررات اختصاصی درباره خسارات ناشی از هوش مصنوعی خودمختار، موجب ابهام در تعیین شخص مسئول، نحوه احراز رابطه سببیت و شیوه جبران خسارت شده است. از منظر سیاست جنایی نیز رویکرد تقنینی ایران هنوز پاسخگوی مخاطرات ناشی از سامانه‌های هوشمند نیست و نیازمند اصلاحات ساختاری در حوزه تقنین، نظارت، مسئولیت مدنی و مدیریت ریسک فناوری است. بر این اساس، پژوهش حاضر پیشنهاد می‌کند که با بهره‌گیری از تجارب اتحادیه اروپا، نظامی ترکیبی مبتنی بر مسئولیت مبتنی بر خطر، بیمه اجباری، صندوق جبران خسارت و نظارت پیشینی بر سامانه‌های هوش مصنوعی در حقوق ایران طراحی شود تا ضمن حمایت مؤثر از زیان‌دیدگان، زمینه توسعه مسئولانه فناوری نیز فراهم گردد. این پژوهش علاوه بر تحلیل تطبیقی مسئولیت مدنی، ارتباط میان سیاست جنایی پیشگیرانه و سازوکارهای جبران خسارت را نیز تبیین کرده و مدلی نسبتاً جامع برای اصلاح نظام حقوقی ایران ارائه می‌دهد.

واژگان کلیدی: هوش مصنوعی خودمختار، مسئولیت مدنی، سیاست جنایی، مسئولیت مبتنی بر خطر، جبران خسارت، حقوق تطبیقی، حقوق ایران، اتحادیه اروپا.

طبقه‌بندی: JEL فقه - حقوق - جزا و جرم شناسی - حقوق بین الملل - حقوق خصوصی

A comparative study of civil liability and criminal policy for damages caused by artificial intelligence autonomy; Emphasizing the rights of Iran, the European Union and the United States

Abstract

With the ability to learn, make independent decisions and change behavior based on new data, autonomous artificial intelligence has created one of the most important challenges of contemporary legal systems in the field of civil responsibility and criminal policy. In such a situation, the attribution of damages caused by the independent decisions of artificial intelligence systems to human persons, whether they are producers, developers, operators or owners, has faced serious difficulties and the traditional bases of civil responsibility based on fault, causation and predictability of losses have lost their previous effectiveness to a large extent. On the other hand, the development of artificial intelligence applications in areas such as transportation, medicine, financial services, security and justice has made the criminal policy of governments face the necessity of designing new preventive, monitoring and compensatory mechanisms. In this regard, the present research was carried out with the aim of explaining the basics of civil responsibility due to the autonomy of artificial intelligence and analyzing the criminal policy of the legal systems of Iran, the European Union and the United States, using a descriptive-analytical method and a comparative approach. The research data have been collected and analyzed through the study of library sources, laws, international documents and authentic Persian and English scientific articles. The findings of the research show that the European Union is trying to compensate for the existing gaps by approving specific regulations in the field of artificial intelligence and developing a risk-based liability system, no-fault liability, mandatory insurance and facilitating the proof of litigation; In contrast, the US legal system still relies primarily on product liability, professional malpractice, and case law rules, although there is also a growing trend toward specialized regulation. In Iranian law, despite the possibility of invoking the general rules of civil liability, including the theory of loss, attribution, warranty and jurisprudence, the lack of specific regulations about damages caused by autonomous artificial intelligence has caused ambiguity in determining the responsible person, the method of ascertaining the relationship of causation and the method of compensation. From the point of view of criminal policy, Iran's legislative approach is still not responsive to the risks caused by smart systems and requires structural reforms in the field of legislation, supervision, civil liability and technology risk management. Based on this, the current research suggests that by using the experiences of the European Union, a hybrid system based on risk-based responsibility, compulsory insurance, compensation fund and prior supervision of artificial intelligence systems in Iranian law should be designed in order to provide effective protection to the victims and the responsible development of technology. In addition to the comparative analysis of civil liability, this research also explains the relationship between preventive criminal policy and compensation mechanisms and provides a relatively comprehensive model for reforming Iran's legal system.

Keywords: Autonomous artificial intelligence, Error ۵۰۰ (Server Error)!!۱۵۰۰. That's an error. There was an error. Please try again later. That's all we know., Criminal policy, Risk-based liability, compensation, Comparative law, Iranian rights, European Union.

متن مقاله

پیچیدگی الگوریتم‌های یادگیری عمیق و پدیده «جعبه سیاه» (Black Box) موجب شده است که در بسیاری از موارد، علت واقعی تصمیم سامانه حتی برای طراحان آن نیز قابل شناسایی نباشد؛ موضوعی که اثبات رابطه سببیت و احراز تقصیر را با دشواری‌های فراوان روبه‌رو ساخته است (کشاورز صفی‌ئی، ۱۴۰۴؛ European Commission, ۲۰۲۲).

هم‌زمان با این تحولات، سیاست جنایی نیز با چالش‌های تازه‌ای مواجه شده است. سیاست جنایی نوین تنها به واکنش کیفری پس از وقوع جرم محدود نمی‌شود، بلکه مجموعه‌ای از تدابیر تقنینی، اجرایی، نظارتی، مدیریتی و پیشگیرانه را برای کنترل مخاطرات اجتماعی در بر می‌گیرد. گسترش هوش مصنوعی خودمختار موجب شده است که دولت‌ها علاوه بر اندیشیدن به شیوه جبران خسارات، به دنبال طراحی نظام‌های ارزیابی ریسک، استانداردهای ایمنی، مسئولیت‌پذیری الگوریتمی، شفافیت تصمیمات هوشمند و سازوکارهای نظارت پیشینی باشند. در همین راستا، اتحادیه اروپا با تصویب قانون هوش مصنوعی (AI Act) و اصلاح قواعد مسئولیت مدنی، تلاش کرده است میان حمایت از نوآوری و تضمین حقوق زیان‌دیدگان تعادل برقرار کند، در حالی که ایالات متحده بیشتر بر توسعه تدریجی قواعد مسئولیت محصول، مسئولیت حرفه‌ای و رویه قضایی تکیه دارد (European Parliament, ۲۰۲۴; European Commission, ۲۰۲۲).

در حقوق ایران نیز اگرچه قانون مسئولیت مدنی مصوب ۱۳۳۹، قواعد اتلاف و تسبیب در قانون مدنی و مبانی فقهی همچون قاعده لاضرر، غرور و ضمان ید امکان جبران بسیاری از خسارات را فراهم می‌کنند، اما این مقررات برای مواجهه با خسارات ناشی از سامانه‌های هوش مصنوعی خودمختار طراحی نشده‌اند. از این رو، در مواردی که زیان ناشی از تصمیم مستقل یک سامانه هوشمند باشد، قواعد موجود پاسخ روشنی درباره تعیین مسئول، نحوه توزیع مسئولیت میان اشخاص متعدد، معیار احراز رابطه سببیت و حدود مسئولیت تولیدکنندگان فناوری ارائه نمی‌کنند. این خلأ تقنینی، به‌ویژه با گسترش استفاده از هوش مصنوعی در نظام بانکی، سلامت، حمل‌ونقل، تجارت الکترونیکی و خدمات عمومی، می‌تواند

تحولات شتابان فناوری‌های دیجیتال در دهه اخیر، به‌ویژه توسعه سامانه‌های هوش مصنوعی مبتنی بر یادگیری ماشین، یادگیری عمیق و مدل‌های مولد، ساختارهای سنتی مسئولیت حقوقی را با چالش‌های بنیادین مواجه کرده است. اگرچه در گذشته هوش مصنوعی صرفاً به‌عنوان ابزاری در اختیار انسان تلقی می‌شد، اما نسل جدید سامانه‌های هوشمند با برخورداری از قابلیت تحلیل داده‌های کلان، یادگیری مستمر، تصمیم‌گیری مستقل و حتی اصلاح رفتار بدون دخالت مستقیم انسان، از مرحله ابزار صرف فراتر رفته و به بازیگری مؤثر در بسیاری از عرصه‌های اقتصادی، اجتماعی، امنیتی و قضایی تبدیل شده است. امروزه خودروهای خودران، سامانه‌های تشخیص پزشکی، الگوریتم‌های معاملاتی بازارهای مالی، پهپادهای هوشمند، ربات‌های صنعتی و سامانه‌های تصمیم‌یار قضایی نمونه‌هایی از فناوری‌هایی هستند که بخش مهمی از تصمیمات خود را به‌صورت خودمختار اتخاذ می‌کنند. این تحول اگرچه مزایای فراوانی در افزایش بهره‌وری، کاهش هزینه‌ها و ارتقای کیفیت خدمات به همراه داشته، اما هم‌زمان خطر بروز خسارات مالی، جانی و حتی نقض حقوق بنیادین اشخاص را نیز افزایش داده است (Russell & Norvig, ۲۰۲۱).

در نظام‌های کلاسیک حقوق خصوصی، مسئولیت مدنی بر پایه ارکانی همچون وجود فعل زیان‌بار، ورود ضرر، رابطه سببیت و تقصیر استوار است. این نظام بر این فرض بنا شده که عامل زیان، شخصی انسانی است که رفتار او قابلیت انتساب داشته و می‌توان تقصیر یا بی‌احتیاطی وی را احراز کرد. با این حال، هنگامی که یک سامانه هوش مصنوعی بدون دستور مستقیم انسان و بر اساس الگوریتم‌های یادگیرنده تصمیمی اتخاذ می‌کند که منجر به ورود خسارت می‌شود، تعیین شخص مسئول با ابهام جدی مواجه خواهد شد. در چنین وضعیتی مشخص نیست که مسئولیت باید متوجه تولیدکننده نرم‌افزار، برنامه‌نویس، مالک سامانه، بهره‌بردار، ارائه‌دهنده داده یا حتی مجموعه‌ای از اشخاص باشد. افزون بر این،

شاخه‌های حقوق، اقتصاد، علوم سیاسی و سیاست‌گذاری عمومی را تحت تأثیر قرار داده است. در سال‌های اخیر، پیشرفت الگوریتم‌های یادگیری ماشین، یادگیری عمیق و شبکه‌های عصبی موجب شده است سامانه‌های هوشمند از مرحله اجرای صرف دستورات برنامه‌نویس عبور کرده و بتوانند با تحلیل حجم عظیمی از داده‌ها، الگوهای جدید را شناسایی، رفتار خود را اصلاح و در بسیاری از موارد بدون مداخله مستقیم انسان تصمیم‌گیری کنند. همین ویژگی، یعنی «خودمختاری»، مهم‌ترین عامل تمایز نسل جدید هوش مصنوعی با نرم‌افزارهای سنتی است و منشأ بسیاری از مسائل نوظهور حقوقی در حوزه مسئولیت مدنی و سیاست جنایی به شمار می‌رود (Russell & Norvig, ۲۰۲۱).

در ادبیات حقوقی، مفهوم خودمختاری هوش مصنوعی به میزان استقلال سامانه در اتخاذ تصمیم، انتخاب میان گزینه‌های مختلف و اجرای اقدامات بدون دخالت لحظه‌ای انسان اشاره دارد. بنابراین، هرچه میزان وابستگی سامانه به دستورات مستقیم انسان کمتر و امکان یادگیری و اصلاح رفتار آن بیشتر باشد، درجه خودمختاری افزایش می‌یابد. از این منظر، خودروهای خودران، ربات‌های جراح، سامانه‌های معاملات مالی الگوریتمی، پهپادهای هوشمند و برخی سامانه‌های امنیت سایبری از نمونه‌های بارز هوش مصنوعی خودمختار محسوب می‌شوند که تصمیمات آن‌ها می‌تواند آثار حقوقی مستقیم بر اشخاص ثالث بر جای گذارد (کشاورز صفی‌ئی، ۱۴۰۴).

ویژگی اساسی این سامانه‌ها، قابلیت یادگیری مستمر از محیط است. برخلاف نرم‌افزارهای سنتی که رفتار آن‌ها صرفاً تابع کدهای اولیه برنامه‌نویسی است، سامانه‌های مبتنی بر یادگیری ماشین با دریافت داده‌های جدید، مدل تصمیم‌گیری خود را تغییر می‌دهند و در نتیجه، خروجی آن‌ها همواره قابل پیش‌بینی نیست. این ویژگی، اگرچه دقت عملکرد را افزایش می‌دهد، اما هم‌زمان تعیین منشأ خطا و انتساب خسارت را با پیچیدگی روبه‌رو می‌کند؛ زیرا ممکن است رفتار زیان‌بار سامانه نه ناشی از نقص اولیه برنامه، بلکه نتیجه فرایند یادگیری پس از بهره‌برداری باشد (European Commission, ۲۰۲۲).

امنیت حقوقی شهروندان و اعتماد عمومی به فناوری‌های نوین را با مخاطره مواجه سازد (کشاورز صفی‌ئی، ۱۴۰۴).

اهمیت این موضوع زمانی دوچندان می‌شود که ملاحظه شود بیشتر پژوهش‌های داخلی، مسئولیت مدنی هوش مصنوعی را صرفاً از منظر قواعد سنتی مسئولیت یا مسئولیت تولیدکننده بررسی کرده‌اند و کمتر به ارتباط میان خودمختاری هوش مصنوعی، تحول مبانی مسئولیت مدنی و سیاست جنایی پیشگیرانه پرداخته‌اند. همچنین، اغلب مطالعات تطبیقی موجود تنها بر یکی از نظام‌های حقوقی متمرکز بوده و تحلیل جامعی از تفاوت‌های رویکرد ایران، اتحادیه اروپا و ایالات متحده در زمینه مدیریت خسارات ناشی از هوش مصنوعی ارائه نکرده‌اند. این خلأ پژوهشی ضرورت انجام مطالعه حاضر را آشکار می‌سازد.

بر همین اساس، هدف اصلی این پژوهش، تحلیل تطبیقی مبانی مسئولیت مدنی و سیاست جنایی در قبال خسارات ناشی از خودمختاری هوش مصنوعی با تأکید بر حقوق ایران، اتحادیه اروپا و ایالات متحده است. در این راستا، پژوهش حاضر می‌کوشد ضمن بررسی کارآمدی مبانی سنتی مسئولیت مدنی، میزان قابلیت انطباق آن‌ها با فناوری‌های نوین را ارزیابی کرده، تحولات تقنینی و قضایی نظام‌های حقوقی پیشرو را تحلیل کند و در نهایت، راهکارهایی برای اصلاح نظام حقوقی ایران و طراحی الگویی متناسب با اقتضات هوش مصنوعی خودمختار ارائه دهد. نوآوری این پژوهش در آن است که برای نخستین بار، رابطه میان مسئولیت مدنی و سیاست جنایی در مواجهه با خسارات ناشی از خودمختاری هوش مصنوعی را در قالب یک مطالعه تطبیقی و با تکیه بر تحولات جدید تقنینی، به‌ویژه مقررات اتحادیه اروپا، بررسی کرده و تلاش دارد الگویی تلفیقی مبتنی بر مسئولیت مبتنی بر خطر، بیمه اجباری، مدیریت ریسک و نظارت پیشینی را برای نظام حقوقی ایران پیشنهاد کند.

۱. مبانی مفهومی و نظری خودمختاری هوش مصنوعی
ظهور هوش مصنوعی یکی از مهم‌ترین تحولات فناورانه قرن بیست و یکم محسوب می‌شود که آثار آن فراتر از حوزه علوم رایانه، بسیاری از

مسئولیت ناشی از هوش مصنوعی تصویب نشده است، اما دادگاه‌ها و نهادهای تنظیم‌گر تلاش کرده‌اند از طریق توسعه قواعد مسئولیت محصول، مسئولیت حرفه‌ای و الزامات ایمنی، بخشی از خلأهای موجود را جبران کنند. در این نظام، همچنان محور اصلی مسئولیت، رفتار اشخاص انسانی است، هرچند افزایش پیچیدگی سامانه‌های هوشمند، ضرورت اصلاح تدریجی قواعد سنتی را آشکار ساخته است (Abbott, ۲۰۲۰).

در حقوق ایران، مفهوم خودمختاری هوش مصنوعی هنوز وارد قوانین موضوعه نشده و مقررات اختصاصی درباره مسئولیت ناشی از عملکرد مستقل سامانه‌های هوشمند وجود ندارد. با این حال، قواعد عام قانون مسئولیت مدنی، مواد مربوط به اتلاف و تسبیب در قانون مدنی و نیز قواعد فقهی همچون لاضرر، غرور، اتلاف و ضمان ید، ظرفیت‌هایی برای جبران خسارات فراهم می‌آورند. با وجود این، این قواعد در شرایطی تدوین شده‌اند که فرض غالب، انتساب مستقیم رفتار زیان‌بار به انسان بوده است؛ بنابراین، در مواجهه با سامانه‌هایی که به‌طور مستقل تصمیم می‌گیرند و رفتار خود را تغییر می‌دهند، پاسخگوی همه مسائل جدید نیستند. از این‌رو، بسیاری از پژوهشگران بر این باورند که حقوق ایران نیز همانند نظام‌های حقوقی پیشرو، نیازمند بازنگری در مبانی مسئولیت مدنی و طراحی قواعد اختصاصی برای مدیریت مخاطرات ناشی از هوش مصنوعی است (کشاورز صفی‌ئی، ۱۴۰۴).

بر این اساس، خودمختاری هوش مصنوعی صرفاً یک ویژگی فنی نیست، بلکه مفهومی حقوقی است که مستقیماً بر ارکان مسئولیت مدنی، معیارهای انتساب خسارت، شیوه‌های جبران زیان و سیاست جنایی دولت‌ها اثر می‌گذارد. هرچه استقلال سامانه‌ها افزایش یابد، فاصله میان قواعد سنتی مسئولیت و واقعیت‌های فناوری بیشتر خواهد شد و این امر ضرورت تدوین چارچوب‌های حقوقی نوین را اجتناب‌ناپذیر می‌سازد. از این‌رو، پیش از بررسی مبانی مسئولیت مدنی، لازم است تأثیر خودمختاری هوش مصنوعی بر ارکان کلاسیک مسئولیت و به‌ویژه بر مفهوم تقصیر، رابطه سببیت و انتساب فعل زیان‌بار مورد تحلیل قرار گیرد؛ موضوعی که در بخش بعدی مقاله به آن پرداخته خواهد شد.

از منظر حقوقی، خودمختاری هوش مصنوعی به معنای استقلال شخصیت حقوقی سامانه نیست. در اغلب نظام‌های حقوقی، هوش مصنوعی هنوز فاقد شخصیت حقوقی مستقل تلقی می‌شود و نمی‌تواند همانند اشخاص حقیقی یا حقوقی، طرف حق و تکلیف قرار گیرد. با وجود این، افزایش استقلال عملکرد سامانه‌ها سبب شده است که نظریه‌های سنتی انتساب مسئولیت، به‌ویژه نظریه تقصیر، با چالش‌های جدی مواجه شوند. در واقع، هر اندازه نقش انسان در تصمیم نهایی کاهش یابد، احراز تقصیر تولیدکننده، طراح یا بهره‌بردار دشوارتر خواهد شد؛ موضوعی که برخی نویسندگان از آن با عنوان «شکاف مسئولیت» (Responsibility Gap) یاد می‌کنند (Matthias, ۲۰۰۴).

در پاسخ به این چالش، رویکردهای متفاوتی در نظام‌های حقوقی شکل گرفته است. گروهی همچنان معتقدند که هوش مصنوعی صرفاً ابزار محسوب می‌شود و مسئولیت باید بر عهده اشخاص انسانی باقی بماند؛ اعم از تولیدکننده، توسعه‌دهنده، مالک یا بهره‌بردار. در مقابل، برخی صاحب‌نظران پیشنهاد کرده‌اند که با توجه به استقلال روزافزون سامانه‌های هوشمند، باید مبانی سنتی مسئولیت مدنی مورد بازنگری قرار گیرد و الگوهایی مانند مسئولیت مبتنی بر خطر، مسئولیت بدون تقصیر، بیمه اجباری و صندوق‌های جبران خسارت جایگزین یا مکمل نظام سنتی شوند (European Parliament, ۲۰۲۰).

حقوق اتحادیه اروپا در سال‌های اخیر بیش از سایر نظام‌های حقوقی به این تحول توجه نشان داده است. تصویب قانون هوش مصنوعی اتحادیه اروپا (AI Act) و پیشنهاد دستورالعمل مسئولیت مدنی هوش مصنوعی (AI Liability Directive) نشان می‌دهد که قانون‌گذار اروپایی، خودمختاری سامانه‌های هوشمند را نه صرفاً یک مسئله فناورانه، بلکه موضوعی مرتبط با امنیت حقوقی، حمایت از مصرف‌کنندگان و حفظ اعتماد عمومی می‌داند. بر همین اساس، مقررات اروپایی بر ارزیابی ریسک، شفافیت الگوریتم‌ها، قابلیت ردیابی تصمیمات و تسهیل اثبات دعوای مسئولیت مدنی تأکید کرده‌اند (European Parliament, ۲۰۲۴).

در ایالات متحده نیز اگرچه تاکنون قانون جامع و یکپارچه‌ای درباره

برای الزام شخص به جبران خسارت کافی نیست، بلکه باید میان رفتار منتسب به وی و زیان ایجاد شده رابطه‌ای مستقیم و قابل اثبات وجود داشته باشد. در سامانه‌های هوش مصنوعی، این رابطه به دلیل تعدد عوامل مؤثر در تصمیم‌گیری، از جمله کیفیت داده‌های آموزشی، طراحی الگوریتم، نحوه به‌روزرسانی نرم‌افزار، مداخله بهره‌بردار و شرایط محیطی، بسیار پیچیده می‌شود. در عمل، ممکن است خسارت نتیجه مجموعه‌ای از عوامل باشد و نتوان آن را صرفاً به یک شخص منتسب کرد. افزون بر این، بسیاری از الگوریتم‌های یادگیری عمیق از ساختاری برخوردارند که دلیل دقیق تصمیمات آن‌ها برای کاربران و حتی طراحان قابل توضیح نیست؛ وضعیتی که از آن با عنوان «جعبه سیاه الگوریتمی» یاد می‌شود و اثبات رابطه سببیت را با دشواری روبه‌رو می‌سازد (European Parliament, ۲۰۲۴).

چالش سوم، تعیین شخص مسئول است. در زنجیره تولید و بهره‌برداری از یک سامانه هوش مصنوعی، اشخاص متعددی نقش دارند؛ از جمله تولیدکننده سخت‌افزار، توسعه‌دهنده نرم‌افزار، تأمین‌کننده داده، شرکت ارائه‌دهنده خدمات ابری، مالک سامانه و کاربر نهایی. در بسیاری از موارد، خسارت حاصل از تعامل مجموعه‌ای از این عوامل است و نمی‌توان تنها یکی از آن‌ها را مسئول شناخت. از این‌رو، برخی نظام‌های حقوقی به سمت پذیرش مسئولیت تضامنی یا تقسیم مسئولیت میان اشخاص دخیل حرکت کرده‌اند تا از تضییع حقوق زیان‌دیدگان جلوگیری شود (Abbott, ۲۰۲۰).

در برابر این چالش‌ها، نظریه‌های جدیدی در حوزه مسئولیت مدنی مطرح شده است. مهم‌ترین این نظریه‌ها، مسئولیت مبتنی بر خطر است. بر اساس این دیدگاه، هر شخصی که فعالیتی با خطر غیرمتعارف ایجاد می‌کند، باید خسارات ناشی از آن را جبران کند؛ حتی اگر مرتکب تقصیر نشده باشد. فلسفه این نظریه آن است که منافع اقتصادی ناشی از بهره‌برداری از فناوری باید با مسئولیت جبران زیان‌های احتمالی آن همراه باشد. از این منظر، تولیدکننده یا بهره‌بردار سامانه هوش مصنوعی که از منافع اقتصادی آن بهره‌مند می‌شود، مناسب‌ترین شخص برای تحمل ریسک ناشی از

۲. مسئولیت مدنی ناشی از خودمختاری هوش مصنوعی

مسئولیت مدنی یکی از مهم‌ترین نهادهای حقوق خصوصی است که هدف اصلی آن، جبران خسارت زیان‌دیده و اعاده تعادل برهم‌خورده میان اشخاص است. در نظام‌های حقوقی مختلف، مسئولیت مدنی بر این اصل استوار است که هر شخصی که در نتیجه فعل یا ترک فعل خود به دیگری زیان وارد کند، مکلف به جبران آن خواهد بود. این نظام طی دهه‌های متمادی بر پایه مفاهیمی همچون تقصیر، رابطه سببیت، قابلیت انتساب رفتار و ورود ضرر شکل گرفته و قواعد آن عمدتاً با فرض نقش مستقیم انسان در ایجاد زیان تدوین شده است. با ظهور سامانه‌های هوش مصنوعی خودمختار، این پیش‌فرض اساسی دستخوش تحول شده است؛ زیرا در بسیاری از موارد، تصمیم زیان‌بار نه مستقیماً توسط انسان، بلکه توسط سامانه‌ای اتخاذ می‌شود که رفتار خود را بر اساس الگوریتم‌های یادگیرنده و داده‌های جدید تغییر داده است. در نتیجه، پرسش اساسی این است که در چنین شرایطی، مسئولیت مدنی بر عهده چه شخص یا اشخاصی قرار می‌گیرد و آیا قواعد سنتی همچنان توان پاسخگویی به این وضعیت را دارند یا خیر (کشاورز صفی‌ئی، ۱۴۰۴).

نخستین چالش، احراز تقصیر است. در نظریه سنتی مسئولیت، تقصیر به معنای نقض یک تکلیف قانونی یا عرفی از سوی شخص مسئول است. در بسیاری از دعاوی ناشی از هوش مصنوعی، زیان‌دیده باید اثبات کند که تولیدکننده، برنامه‌نویس، توسعه‌دهنده یا بهره‌بردار مرتکب بی‌احتیاطی یا قصور شده است. با این حال، هنگامی که سامانه پس از استقرار و بهره‌برداری، از طریق یادگیری ماشین رفتار خود را تغییر می‌دهد و تصمیمی اتخاذ می‌کند که حتی برای طراح آن نیز قابل پیش‌بینی نبوده است، اثبات تقصیر بسیار دشوار می‌شود. این وضعیت سبب شده است که برخی نویسندگان از ناکارآمدی نظریه تقصیر در مواجهه با سامانه‌های خودمختار سخن بگویند و بر ضرورت حرکت به سوی مبانی نوین مسئولیت تأکید کنند (European Commission, ۲۰۲۲).

دومین چالش، رابطه سببیت است. در مسئولیت مدنی، صرف ورود ضرر

مصنوعی خواهد شد و میان توسعه فناوری و حمایت از حقوق اشخاص تعادل برقرار می‌کند.

۳. سیاست جنایی در قبال خسارات ناشی از خودمختاری هوش مصنوعی گسترش فناوری‌های مبتنی بر هوش مصنوعی، صرفاً مسئله‌ای فنی یا اقتصادی نیست، بلکه به‌طور مستقیم بر امنیت عمومی، حقوق بنیادین اشخاص، نظم اجتماعی و اعتماد عمومی به نظام حقوقی اثر می‌گذارد. از این‌رو، مواجهه با مخاطرات ناشی از هوش مصنوعی دیگر صرفاً در چارچوب مسئولیت مدنی قابل تحلیل نیست، بلکه نیازمند اتخاذ سیاست جنایی جامع است. سیاست جنایی در مفهوم نوین خود، مجموعه‌ای از تدابیر تقنینی، قضایی، اجرایی، اداری، فرهنگی و فناورانه است که با هدف پیشگیری از وقوع رفتارهای زیان‌بار، کاهش مخاطرات اجتماعی، حمایت از بزه‌دیدگان و تضمین اجرای عدالت اتخاذ می‌شود. بر این اساس، خسارات ناشی از خودمختاری هوش مصنوعی نیز باید در قالب سیاست جنایی پیشگیرانه و حمایتی مورد بررسی قرار گیرد، نه صرفاً از منظر پاسخ‌های سنتی پس از وقوع زیان (دلماس مارتی، ۱۳۸۱؛ نجفی ابرندآبادی، ۱۳۹۷).

یکی از مهم‌ترین ویژگی‌های سیاست جنایی در حوزه هوش مصنوعی، تغییر رویکرد از واکنش پسینی به مدیریت پیشینی خطر است. در نظام‌های سنتی، قانون‌گذار غالباً پس از وقوع خسارت، از طریق مسئولیت مدنی یا ضمانت اجرای کیفری واکنش نشان می‌دهد؛ اما در فناوری‌های هوشمند، به دلیل احتمال بروز خسارات گسترده و گاه غیرقابل جبران، سیاست جنایی به سمت شناسایی و کنترل خطر پیش از وقوع زیان حرکت کرده است. این رویکرد بر ارزیابی ریسک، استانداردسازی سامانه‌ها، نظارت مستمر، شفافیت الگوریتمی، ممیزی فنی و مسئولیت‌پذیری توسعه‌دهندگان استوار است و هدف آن، کاهش احتمال وقوع خسارت پیش از ورود زیان به اشخاص است (European Parliament, ۲۰۲۴).

در این چارچوب، اصل پیشگیری به یکی از مهم‌ترین اصول سیاست جنایی

عملکرد آن نیز خواهد بود (کشاورز صفی‌ئی، ۱۴۰۴).

در همین راستا، اتحادیه اروپا با اصلاح مقررات مسئولیت مدنی و ارائه چارچوب‌های اختصاصی برای هوش مصنوعی، تلاش کرده است بار اثبات را در برخی موارد به نفع زیان‌دیده تسهیل کند. در مواردی که بهره‌بردار یا تولیدکننده الزامات قانونی مربوط به شفافیت، ثبت اطلاعات یا مدیریت ریسک را رعایت نکرده باشد، امکان ایجاد اماره قانونی درباره رابطه سببیت و مسئولیت پیش‌بینی شده است. این رویکرد نشان می‌دهد که قانون‌گذار اروپایی به جای ایجاد شخصیت حقوقی مستقل برای هوش مصنوعی، مسئولیت را همچنان متوجه اشخاص انسانی می‌داند، اما قواعد اثبات و انتساب را متناسب با ویژگی‌های فناوری اصلاح کرده است (European Commission, ۲۰۲۲).

در حقوق ایران نیز هرچند مقررات اختصاصی در این زمینه وجود ندارد، اما برخی ظرفیت‌های حقوقی قابل استفاده است. قواعد اتلاف و تسبیب در قانون مدنی، قانون مسئولیت مدنی مصوب ۱۳۳۹ و همچنین قواعد فقهی مانند لاضرر، غرور و اتلاف می‌توانند مبنای الزام به جبران خسارت باشند. با این حال، این قواعد برای شرایطی تدوین شده‌اند که عامل زیان رفتاری انسانی داشته است و درباره خسارات ناشی از تصمیم مستقل یک سامانه هوشمند حکم صریحی ارائه نمی‌کنند. از این‌رو، در فرضی که خودرو خودران، سامانه پزشکی یا الگوریتم مالی بدون دخالت مستقیم انسان موجب خسارت شود، دادگاه با ابهام در تعیین شخص مسئول و نحوه احراز رابطه سببیت مواجه خواهد شد. این خلأ، امنیت حقوقی را کاهش داده و می‌تواند مانعی برای توسعه مسئولانه فناوری نیز باشد.

از منظر نگارنده، راهکار مناسب برای حقوق ایران، تکیه صرف بر نظریه تقصیر نیست، بلکه باید نظامی ترکیبی طراحی شود که در آن، اصل بر مسئولیت مبتنی بر خطر در فعالیت‌های پرریسک هوش مصنوعی باشد و در کنار آن، بیمه اجباری مسئولیت، صندوق جبران خسارت، الزامات ثبت و نگهداری داده‌های عملکرد سامانه و استانداردهای فنی نیز پیش‌بینی شود. چنین الگویی ضمن حمایت مؤثر از زیان‌دیدگان، مانع از تحمیل مسئولیت نامحدود و غیرمنصفانه بر نوآوران و سرمایه‌گذاران حوزه هوش

سببیت ایجاد شود تا زیان دیدگان بدون تحمل بار سنگین اثبات، بتوانند به حقوق خود دست یابند. چنین تدابیری، علاوه بر حمایت از اشخاص، اعتماد عمومی به توسعه فناوری‌های نوین را نیز افزایش می‌دهد (کشاورز صفی‌ئی، ۱۴۰۴).

در حقوق ایران، مفهوم سیاست جنایی در حوزه هوش مصنوعی هنوز در مراحل ابتدایی قرار دارد. اگرچه در اسناد کلان کشور، توسعه فناوری‌های نوین مورد توجه قرار گرفته است، اما تاکنون قانون جامع یا سیاست تقنینی مشخصی برای مدیریت مخاطرات ناشی از سامانه‌های خودمختار تدوین نشده است. همچنین، مقررات موجود عمدتاً پس از وقوع زیان قابلیت اجرا دارند و سازوکارهای پیشگیرانه مانند ارزیابی ریسک، نظارت بر الگوریتم‌ها، ثبت عملکرد سامانه‌های هوشمند یا الزام به ممیزی مستقل در قوانین ایران پیش‌بینی نشده است. این خلأ، نه تنها از حقوق زیان دیدگان حمایت کافی نمی‌کند، بلکه امنیت حقوقی سرمایه‌گذاران و فعالان حوزه فناوری را نیز با ابهام مواجه می‌سازد.

از منظر سیاست جنایی قضایی نیز، دادگاه‌های ایران در صورت بروز خسارات ناشی از هوش مصنوعی ناگزیر به استناد قواعد عام مسئولیت مدنی و اصول کلی فقهی خواهند بود. این امر ممکن است به صدور آرای متفاوت در پرونده‌های مشابه و کاهش وحدت رویه منجر شود. در مقابل، وجود چارچوب‌های قانونی اختصاصی، ضمن افزایش پیش‌بینی‌پذیری تصمیمات قضایی، امکان توزیع عادلانه مسئولیت میان اشخاص دخیل در طراحی، تولید و بهره‌برداری از سامانه‌های هوشمند را فراهم می‌آورد.

در مجموع، بررسی تطبیقی نشان می‌دهد که سیاست جنایی مؤثر در حوزه هوش مصنوعی باید بر چهار محور اساسی استوار باشد: پیشگیری مبتنی بر ارزیابی ریسک، شفافیت و قابلیت توضیح تصمیمات الگوریتمی، حمایت مؤثر از زیان دیدگان و مسئولیت‌پذیری اشخاص دخیل در چرخه تولید و بهره‌برداری از سامانه‌های هوشمند. تحقق این اهداف مستلزم آن است که سیاست جنایی از رویکرد سنتی واکنش پس از وقوع خسارت فاصله گرفته و به سمت مدیریت فعال مخاطرات فناوری حرکت کند. در این میان، تجربه اتحادیه اروپا می‌تواند الگوی مناسبی برای اصلاح نظام

تبدیل شده است. بر اساس این اصل، هرچه احتمال ورود خسارت ناشی از یک فناوری بیشتر باشد، میزان نظارت و الزامات قانونی نیز باید افزایش یابد. به همین دلیل، اتحادیه اروپا در قانون هوش مصنوعی، سامانه‌ها را بر اساس میزان خطر به چهار دسته تقسیم کرده و برای هر گروه، الزامات متفاوتی مقرر کرده است. سامانه‌هایی که در حوزه‌هایی مانند سلامت، حمل‌ونقل، اجرای قانون، زیرساخت‌های حیاتی و عدالت قضایی به کار گرفته می‌شوند، در زمره سامانه‌های «پرخطر» قرار گرفته و مشمول الزامات سخت‌گیرانه‌ای همچون ارزیابی انطباق، مستندسازی فنی، ثبت سوابق عملکرد، نظارت انسانی و مدیریت مستمر ریسک هستند. این رویکرد نشان می‌دهد که سیاست جنایی اتحادیه اروپا بیش از آنکه بر مجازات پس از وقوع خسارت تأکید کند، بر پیشگیری ساختاری از وقوع خطر متمرکز است (European Parliament, ۲۰۲۴).

افزون بر پیشگیری، شفافیت و قابلیت توضیح تصمیمات هوش مصنوعی نیز از ارکان سیاست جنایی نوین به شمار می‌رود. بسیاری از خسارات ناشی از سامانه‌های هوشمند، نتیجه تصمیماتی است که منطق آن‌ها برای کاربران، ناظران و حتی توسعه‌دهندگان قابل درک نیست. این وضعیت، نه تنها اثبات مسئولیت مدنی را دشوار می‌کند، بلکه اعتماد عمومی به فناوری را نیز کاهش می‌دهد. از این‌رو، مقررات جدید اروپایی توسعه‌دهندگان را مکلف کرده است تا اطلاعات لازم درباره نحوه عملکرد سامانه، منابع داده، محدودیت‌های الگوریتم و شیوه تصمیم‌گیری را ثبت و نگهداری کنند تا در صورت بروز خسارت، امکان بررسی و ارزیابی قضایی فراهم باشد (European Commission, ۲۰۲۲).

یکی دیگر از ابعاد سیاست جنایی، حمایت مؤثر از زیان دیدگان است. در بسیاری از موارد، اثبات تقصیر یا رابطه سببیت برای اشخاص زیان دیده بسیار دشوار و پرهزینه است. بنابراین، سیاست جنایی کارآمد باید علاوه بر تعیین مسئول، سازوکارهایی را برای تسهیل جبران خسارت پیش‌بینی کند. به همین دلیل، در برخی نظام‌های حقوقی پیشنهاد شده است که برای فعالیتهای پرریسک مبتنی بر هوش مصنوعی، بیمه اجباری مسئولیت، صندوق‌های جبران خسارت و اماره‌های قانونی در اثبات رابطه

خودمختار، تصمیم نهایی ممکن است نتیجه فرایند یادگیری ماشین و بدون دخالت مستقیم انسان باشد. دوم آنکه تعیین رابطه سببیت میان رفتار تولیدکننده، برنامه‌نویس، مالک یا بهره‌بردار و خسارت ایجادشده، به دلیل پیچیدگی الگوریتم‌ها و تعدد عوامل مؤثر، با ابهام فراوان روبه‌رو است. سوم آنکه در حقوق ایران مقرراتی درباره الزام به ثبت عملکرد سامانه، شفافیت الگوریتمی، بیمه اجباری یا صندوق جبران خسارت وجود ندارد؛ در نتیجه، زیان‌دیده ناگزیر است بار اثبات دعوا را مطابق قواعد عمومی بر عهده گیرد که این امر در بسیاری از موارد، امکان جبران مؤثر خسارت را کاهش می‌دهد (کشاورز صفی‌ئی، ۱۴۰۴).

از منظر سیاست جنایی نیز، رویکرد حقوق ایران عمدتاً واکنشی است؛ یعنی قانون‌گذار پس از وقوع خسارت، امکان مراجعه به قواعد عام مسئولیت مدنی را فراهم کرده است، اما سازوکارهای پیشگیرانه برای کنترل مخاطرات فناوری‌های هوشمند هنوز توسعه نیافته‌اند. این وضعیت نشان می‌دهد که نظام حقوقی ایران بیش از هر چیز نیازمند تدوین مقررات اختصاصی درباره مدیریت ریسک، مسئولیت تولیدکنندگان و حمایت از زیان‌دیدگان است.

۲-۴. اتحادیه اروپا

اتحادیه اروپا در سال‌های اخیر پیشرفته‌ترین چارچوب حقوقی را در زمینه تنظیم فعالیت سامانه‌های هوش مصنوعی ایجاد کرده است. تصویب قانون هوش مصنوعی (AI Act) و ارائه پیشنهاد دستورالعمل مسئولیت مدنی هوش مصنوعی (AI Liability Directive) بیانگر آن است که قانون‌گذار اروپایی، هوش مصنوعی را صرفاً موضوعی فناورانه نمی‌داند، بلکه آن را مسئله‌ای مرتبط با حقوق بنیادین، حمایت از مصرف‌کنندگان، رقابت سالم و امنیت عمومی تلقی می‌کند (European Parliament, ۲۰۲۴).

برخلاف حقوق ایران، رویکرد اتحادیه اروپا بر مدیریت پیشینی خطر استوار است. بر اساس قانون هوش مصنوعی، سامانه‌ها بر مبنای سطح خطر طبقه‌بندی شده‌اند و هرچه احتمال ورود خسارت بیشتر باشد، الزامات قانونی نیز شدیدتر خواهد بود. برای سامانه‌های پرخطر، رعایت الزامات

حقوقی ایران باشد، مشروط بر آنکه این الگو با مبانی حقوقی و فقهی ایران سازگار شده و متناسب با ساختار تقنینی کشور بومی‌سازی شود.

۴. مطالعه تطبیقی مسئولیت مدنی و سیاست جنایی در قبال خسارات ناشی از خودمختاری هوش مصنوعی در حقوق ایران، اتحادیه اروپا و ایالات متحده

گسترش کاربرد سامانه‌های هوش مصنوعی خودمختار، نظام‌های حقوقی را ناگزیر ساخته است تا در قواعد سنتی مسئولیت مدنی و سیاست جنایی بازنگری کنند. با وجود اشتراک نظام‌های حقوقی در هدف حمایت از زیان‌دیدگان و حفظ امنیت عمومی، شیوه تحقق این اهداف در ایران، اتحادیه اروپا و ایالات متحده متفاوت است. این تفاوت‌ها بیش از هر چیز ناشی از اختلاف در مبانی حقوقی، میزان توسعه فناوری، سیاست‌های تقنینی و نقش دولت در تنظیم‌گری فناوری‌های نوظهور است. بررسی تطبیقی این سه نظام حقوقی می‌تواند ضمن آشکار ساختن نقاط قوت و ضعف هر یک، زمینه ارائه الگویی متناسب با نظام حقوقی ایران را فراهم کند (European Commission, ۲۰۲۲).

۱-۴. حقوق ایران

در نظام حقوقی ایران تاکنون قانون اختصاصی درباره مسئولیت مدنی ناشی از عملکرد سامانه‌های هوش مصنوعی خودمختار به تصویب نرسیده است. از این‌رو، دعاوی احتمالی ناشی از خسارات این سامانه‌ها باید بر اساس قواعد عام مسئولیت مدنی، قانون مدنی و اصول فقهی مورد رسیدگی قرار گیرند. مهم‌ترین مبانی قابل استناد عبارت‌اند از قانون مسئولیت مدنی مصوب ۱۳۳۹، قواعد اتلاف و تسبیب مندرج در قانون مدنی و قواعد فقهی همچون لاضرر، اتلاف، تسبیب و غرور (صفایی، ۱۳۹۹).

با وجود ظرفیت این قواعد، اجرای آن‌ها در حوزه هوش مصنوعی با دشواری‌هایی همراه است. نخست آنکه اغلب این مقررات بر فرض انتساب مستقیم رفتار زیان‌بار به انسان استوارند، در حالی که در سامانه‌های

این حال، هرچه سامانه از استقلال بیشتری برخوردار باشد، اثبات تقصیر و رابطه سببیت نیز دشوارتر خواهد شد؛ موضوعی که در سال‌های اخیر زمینه طرح پیشنهاد‌های مختلف برای اصلاح قوانین را فراهم کرده است (Abbott, ۲۰۲۰).

از منظر سیاست جنایی، ایالات متحده بیش از آنکه بر مقررات پیشگیرانه جامع تکیه کند، به انعطاف‌پذیری رویه قضایی، خودتنظیمی شرکت‌های فناوری و تدوین استانداردهای تخصصی توسط نهادهای اجرایی توجه دارد. این رویکرد اگرچه نوآوری را تسهیل می‌کند، اما ممکن است به دلیل نبود قواعد یکپارچه، موجب ناهمگونی در حمایت از زیان‌دیدگان شود.

۴-۴. تحلیل تطبیقی

مقایسه سه نظام حقوقی نشان می‌دهد که حقوق ایران همچنان در مرحله اتکا به قواعد سنتی مسئولیت مدنی قرار دارد و فاقد چارچوب اختصاصی برای مدیریت مخاطرات ناشی از هوش مصنوعی است. در مقابل، اتحادیه اروپا با اتخاذ رویکردی آینده‌نگر، تلاش کرده است از طریق تلفیق مسئولیت مدنی، مدیریت ریسک و سیاست جنایی پیشگیرانه، میان حمایت از نوآوری و حمایت از زیان‌دیدگان تعادل برقرار کند. ایالات متحده نیز اگرچه از انعطاف رویه قضایی بهره می‌برد، اما نبود قانون جامع، پاسخگویی یکسان به چالش‌های ناشی از خودمختاری هوش مصنوعی را دشوار ساخته است.

این مطالعه تطبیقی نشان می‌دهد که اصلاح نظام حقوقی ایران نباید صرفاً به تصویب مقرراتی درباره مسئولیت مدنی محدود شود، بلکه لازم است سیاستی جامع اتخاذ شود که در آن، الزامات ارزیابی ریسک، شفافیت الگوریتمی، مسئولیت مبتنی بر خطر در فعالیت‌های پرریسک، بیمه اجباری، ثبت و نگهداری داده‌های عملکرد سامانه و ایجاد صندوق جبران خسارت به‌صورت هماهنگ پیش‌بینی شود. تنها در چنین صورتی می‌توان ضمن حفظ امنیت حقوقی اشخاص، زمینه توسعه مسئولانه فناوری‌های هوش مصنوعی را نیز فراهم کرد.

مربوط به ارزیابی ریسک، کیفیت داده‌های آموزشی، نظارت انسانی، ثبت اطلاعات، شفافیت و مستندسازی فنی الزامی است. عدم رعایت این الزامات می‌تواند در دعاوی مسئولیت مدنی نیز به زیان تولیدکننده یا بهره‌بردار مورد استناد قرار گیرد (European Parliament, ۲۰۲۴). در حوزه مسئولیت مدنی، اتحادیه اروپا به جای ایجاد شخصیت حقوقی مستقل برای هوش مصنوعی، همچنان مسئولیت را متوجه اشخاص انسانی می‌داند، اما برای حمایت از زیان‌دیدگان، قواعد اثبات دعوا را اصلاح کرده است. به‌ویژه در مواردی که تولیدکننده یا بهره‌بردار الزامات قانونی را رعایت نکرده باشد، دادگاه می‌تواند اماره‌ای به نفع زیان‌دیده درباره رابطه سببیت ایجاد کند. این نوآوری موجب کاهش دشواری اثبات دعوا و افزایش حمایت از اشخاص زیان‌دیده می‌شود (European Commission, ۲۰۲۲).

سیاست جنایی اتحادیه اروپا نیز بر اصل پیشگیری، مسئولیت‌پذیری فناورانه، شفافیت الگوریتمی و حفاظت از حقوق بنیادین استوار است. در این نظام، مسئولیت مدنی بخشی از سیاست جامع تنظیم‌گری فناوری محسوب می‌شود و با ابزارهای نظارتی، استانداردهای فنی و ضمانت اجرای اداری تکمیل می‌شود.

۴-۳. ایالات متحده آمریکا

برخلاف اتحادیه اروپا، ایالات متحده تاکنون قانون جامع فدرالی درباره مسئولیت مدنی ناشی از هوش مصنوعی تصویب نکرده است. رویکرد آمریکا بیشتر مبتنی بر توسعه تدریجی قواعد موجود از طریق رویه قضایی، مسئولیت محصول، مسئولیت حرفه‌ای و مقررات بخشی است (Abbott, ۲۰۲۰).

در دعاوی مرتبط با هوش مصنوعی، دادگاه‌های آمریکا معمولاً تلاش می‌کنند مسئولیت را بر اساس قواعد سنتی حقوق مسئولیت مدنی تعیین کنند. اگر خسارت ناشی از نقص طراحی یا تولید باشد، قواعد مسئولیت محصول اعمال می‌شود و اگر بهره‌بردار در استفاده از سامانه مرتکب بی‌احتیاطی شده باشد، مسئولیت وی بر مبنای تقصیر بررسی می‌شود. با

شخصی که رفتار وی قابل سرزنش باشد کاهش می‌یابد. به همین دلیل، برخی نظام‌های حقوقی به سمت پذیرش معیارهای عینی‌تر مسئولیت حرکت کرده‌اند.

۵-۲. ضرورت پذیرش مسئولیت مبتنی بر خطر یکی از راهکارهای مطرح برای جبران کاستی‌های نظریه تقصیر، پذیرش مسئولیت مبتنی بر خطر است. بر اساس این نظریه، شخصی که فعالیتی خطرآفرین را ایجاد و از منافع اقتصادی آن بهره‌مند می‌شود، باید مسئول پیامدهای زیان‌بار آن نیز باشد؛ حتی اگر مرتکب تقصیر نشده باشد.

هوش مصنوعی در برخی حوزه‌ها مانند پزشکی، حمل‌ونقل، امنیت و امور مالی دارای ریسک ذاتی است. بنابراین، می‌توان استدلال کرد که بهره‌برداران و تولیدکنندگان این فناوری‌ها باید بخشی از خطر ایجادشده را تحمل کنند. این رویکرد علاوه بر حمایت از زیان‌دیدگان، انگیزه بیشتری برای طراحی ایمن‌تر، کنترل کیفیت داده‌ها و نظارت مستمر بر عملکرد سامانه ایجاد می‌کند.

البته مسئولیت مبتنی بر خطر نباید به‌گونه‌ای اعمال شود که مانع توسعه فناوری گردد. از این‌رو، مناسب‌ترین راهکار، پذیرش مسئولیت مطلق و نامحدود نیست، بلکه ایجاد نظامی متعادل است که میان حمایت از حقوق اشخاص و تشویق نوآوری تعادل برقرار کند. در این چارچوب، می‌توان مسئولیت شدیدتر را برای سامانه‌های پرخطر و مسئولیت مبتنی بر تقصیر را برای کاربردهای کم‌خطر پیش‌بینی کرد (European Commission, ۲۰۲۲).

۵-۳. نقش بیمه و صندوق‌های جبران خسارت یکی دیگر از راهکارهای پیشنهادی برای مدیریت خسارات ناشی از هوش مصنوعی، ایجاد سازوکارهای جبرانی مستقل از اثبات تقصیر است. تجربه نظام‌های حقوقی مختلف نشان داده است که در حوزه‌های دارای ریسک اجتماعی بالا، مانند حوادث رانندگی یا خسارات پزشکی، بیمه اجباری و صندوق‌های جبران خسارت می‌توانند نقش مؤثری در حمایت از

۵. تحلیل چالش‌های انتساب مسئولیت و ارائه الگوی مطلوب برای حقوق ایران یکی از اساسی‌ترین مسائل حقوقی ناشی از خودمختاری هوش مصنوعی، تعیین مبنای انتساب مسئولیت است. در نظام سنتی مسئولیت مدنی، فعل زیان‌بار معمولاً به شخص معینی قابل انتساب است؛ زیرا رفتار انسانی، مبتنی بر اراده، آگاهی و امکان کنترل است. اما در سامانه‌های هوش مصنوعی خودمختار، تصمیم‌نهایی ممکن است حاصل تعامل پیچیده میان الگوریتم، داده‌های آموزشی، محیط عملیاتی و فرایند یادگیری سیستم باشد. در چنین شرایطی، مفهوم سنتی «عامل زیان» با ابهام مواجه شده و ضرورت بازتعریف معیارهای انتساب مسئولیت آشکار می‌شود (Matthias, ۲۰۰۴).

۵-۱. ناکارآمدی نظریه تقصیر در مواجهه با هوش مصنوعی خودمختار نظریه تقصیر، مهم‌ترین مبنای مسئولیت مدنی در بسیاری از نظام‌های حقوقی معاصر است. بر اساس این نظریه، شخص زمانی مسئول جبران خسارت است که رفتار وی برخلاف استاندارد متعارف احتیاط بوده و بتوان میان این رفتار و زیان ایجادشده رابطه سببیت برقرار کرد. این نظریه در روابط انسانی سنتی کارآمد است، زیرا امکان ارزیابی رفتار شخص و مقایسه آن با معیارهای متعارف وجود دارد.

با این حال، در حوزه هوش مصنوعی، اثبات تقصیر با موانع جدی روبه‌رو است. برای مثال، اگر یک خودرو خودران به دلیل تصمیم اشتباه الگوریتمی موجب تصادف شود، مشخص نیست آیا باید این خطا را ناشی از تقصیر برنامه‌نویس، نقص طراحی، داده‌های نامناسب آموزشی یا شرایط پیش‌بینی‌نشده محیطی دانست. در بسیاری از موارد، حتی تولیدکننده نیز نمی‌تواند دقیقاً علت تصمیم سامانه را توضیح دهد. بنابراین، الزام زیان‌دیده به اثبات تقصیر شخص مشخص، ممکن است عملاً حق جبران خسارت را غیرممکن سازد (Wagner, ۲۰۱۸).

از این منظر، ادامه اتکای مطلق به نظریه تقصیر می‌تواند موجب ایجاد خلأ مسئولیت شود؛ زیرا هرچه استقلال سامانه بیشتر شود، امکان یافتن

نهایی بر عهده اشخاص انسانی مرتبط با طراحی، تولید، کنترل و بهره‌برداری از آن باقی بماند (European Parliament, ۲۰۲۰).

۵-۵. الگوی پیشنهادی برای حقوق ایران با توجه به ویژگی‌های نظام حقوقی ایران و تجربه نظام‌های تطبیقی، به نظر می‌رسد الگوی مناسب برای مواجهه با خسارات ناشی از هوش مصنوعی باید ترکیبی از قواعد سنتی و سازوکارهای نوین باشد. این الگو می‌تواند شامل عناصر زیر باشد:

نخست، تدوین قانون خاص مسئولیت مدنی هوش مصنوعی که ضمن حفظ امکان استناد به قواعد عمومی مسئولیت، احکام ویژه‌ای درباره سامانه‌های خودمختار، اشخاص مسئول، نحوه اثبات دعوا و حدود مسئولیت تعیین کند.

دوم، طبقه‌بندی سامانه‌های هوش مصنوعی بر اساس میزان خطر؛ به گونه‌ای که سامانه‌های مورد استفاده در حوزه‌های حساس مانند سلامت، حمل‌ونقل و امنیت، تحت نظارت شدیدتر قرار گیرند.

سوم، پذیرش مسئولیت مبتنی بر خطر برای فعالیت‌های پرریسک و کاهش بار اثباتی زیان‌دیدگان.

چهارم، الزام به شفافیت الگوریتمی و ثبت سوابق عملکرد سامانه‌ها تا امکان بررسی قضایی و کارشناسی در زمان وقوع خسارت فراهم شود.

پنجم، ایجاد نظام بیمه اجباری و صندوق جبران خسارت برای تضمین جبران سریع و مؤثر زیان‌ها.

ششم، توسعه سیاست جنایی پیشگیرانه از طریق نظارت نهادی، تدوین استانداردهای فنی، آموزش قضات و کارشناسان و ایجاد نهاد تخصصی نظارت بر فناوری‌های هوشمند.

بر این اساس، حقوق ایران می‌تواند بدون پذیرش شخصیت حقوقی مستقل برای هوش مصنوعی، با اصلاح قواعد مسئولیت مدنی و توسعه سیاست جنایی پیشگیرانه، چارچوبی ایجاد کند که هم از حقوق اشخاص حمایت کند و هم مانع توسعه فناوری‌های نوین نشود.

زیان‌دیدگان ایفا کنند.

در زمینه هوش مصنوعی نیز می‌توان برای سامانه‌های پرخطر، الزام به بیمه مسئولیت مدنی را پیش‌بینی کرد. برای مثال، شرکت تولیدکننده خودروهای خودران یا ارائه‌دهنده سامانه پزشکی هوشمند، پیش از عرضه محصول، ملزم به تأمین پوشش بیمه‌ای مناسب شود. همچنین، در مواردی که تعیین مسئول امکان‌پذیر نیست یا خسارت ناشی از نقص ناشناخته فناوری باشد، صندوق جبران خسارت می‌تواند به‌عنوان سازوکار حمایتی وارد عمل شود.

چنین الگویی ضمن جلوگیری از بی‌حقوق ماندن زیان‌دیدگان، از تحمیل مسئولیت غیرمنصفانه بر یک شخص خاص نیز جلوگیری می‌کند و با ماهیت پیچیده فناوری‌های هوشمند سازگارتر است.

۴-۵. ضرورت ایجاد شخصیت حقوقی مستقل برای هوش مصنوعی؛ نقد و بررسی

یکی از مباحث مطرح در حقوق فناوری، اعطای شخصیت حقوقی محدود به برخی سامانه‌های هوش مصنوعی است. طرفداران این دیدگاه معتقدند همان‌گونه که شرکت‌ها دارای شخصیت حقوقی مستقل هستند، برخی سامانه‌های بسیار پیشرفته هوش مصنوعی نیز می‌توانند دارای شخصیت الکترونیکی محدود باشند تا مسئولیت‌های ناشی از عملکرد خود را تحمل کنند.

با وجود جذابیت نظری این دیدگاه، پذیرش آن با انتقادات جدی مواجه است. مهم‌ترین ایراد آن است که هوش مصنوعی فاقد اراده، آگاهی اخلاقی و قصد حقوقی مشابه انسان است و اعطای شخصیت مستقل ممکن است موجب انتقال مسئولیت از اشخاص واقعی به یک موجودیت غیرقابل دسترس شود. همچنین، ایجاد شخصیت حقوقی برای هوش مصنوعی می‌تواند موجب کاهش انگیزه تولیدکنندگان برای رعایت استانداردهای ایمنی شود.

از این‌رو، رویکرد غالب در اتحادیه اروپا و بسیاری از نظام‌های حقوقی این است که هوش مصنوعی همچنان به‌عنوان ابزار تلقی شود و مسئولیت

۶. نتیجه‌گیری و پیشنهادهای تقنینی

تحول شتابان فناوری‌های هوش مصنوعی و حرکت این فناوری‌ها به سمت خودمختاری تصمیم‌گیری، مفاهیم بنیادین حقوق مسئولیت مدنی و سیاست جنایی را با چالش‌هایی اساسی مواجه ساخته است. نظام‌های سنتی مسئولیت مدنی که بر محور رفتار انسانی، تقصیر، قابلیت پیش‌بینی و رابطه سببیت مستقیم شکل گرفته‌اند، در مواجهه با سامانه‌هایی که قادر به یادگیری، اصلاح رفتار و اتخاذ تصمیم مستقل هستند، با محدودیت‌های جدی روبه‌رو شده‌اند. مسئله اصلی دیگر صرفاً تعیین وجود یا عدم وجود تقصیر نیست، بلکه تعیین شیوه‌ای عادلانه برای توزیع خطر میان تولیدکنندگان، توسعه‌دهندگان، بهره‌برداران و زیان‌دیدگان است.

بررسی تطبیقی حقوق ایران، اتحادیه اروپا و ایالات متحده نشان داد که هر یک از این نظام‌ها با توجه به مبانی حقوقی و سیاست‌گذاری خود، رویکرد متفاوتی نسبت به این چالش اتخاذ کرده‌اند. در حقوق ایران، اگرچه قواعد عمومی مسئولیت مدنی، به‌ویژه مقررات قانون مسئولیت مدنی، قواعد اتلاف و تسبیب و مبانی فقهی ظرفیت‌هایی برای جبران خسارت فراهم می‌کنند، اما این قواعد برای شرایطی طراحی شده‌اند که عامل زیان، شخص انسانی با رفتار قابل انتساب باشد. در نتیجه، در برابر خسارات ناشی از تصمیمات مستقل سامانه‌های هوش مصنوعی، مسائل مهمی مانند تعیین شخص مسئول، اثبات رابطه سببیت و ارزیابی تقصیر همچنان بدون پاسخ روشن باقی مانده است.

در مقابل، اتحادیه اروپا با رویکردی پیشگیرانه و مبتنی بر مدیریت خطر، تلاش کرده است چارچوبی جامع برای کنترل مخاطرات هوش مصنوعی ایجاد کند. این نظام ضمن رد ایده اعطای شخصیت حقوقی مستقل به هوش مصنوعی، مسئولیت را همچنان متوجه اشخاص انسانی مرتبط با چرخه طراحی، تولید و بهره‌برداری از سامانه می‌داند، اما با اصلاح قواعد اثبات، ایجاد الزامات شفافیت، ارزیابی ریسک، ثبت عملکرد سامانه‌ها و طبقه‌بندی فناوری‌ها بر اساس سطح خطر، حمایت مؤثرتری از زیان‌دیدگان فراهم کرده است (European Parliament, ۲۰۲۴).

مطالعه نظام حقوقی ایالات متحده نیز نشان داد که این کشور بیشتر بر

انعطاف‌پذیری قواعد سنتی مسئولیت محصول، مسئولیت حرفه‌ای و نقش دادگاه‌ها در توسعه حقوق تأکید دارد. این رویکرد اگرچه زمینه رشد سریع فناوری و نوآوری را فراهم کرده است، اما به دلیل فقدان چارچوب جامع و یکپارچه، ممکن است در حمایت هماهنگ از زیان‌دیدگان با کاستی‌هایی مواجه شود (Abbott, ۲۰۲۰).

یافته اصلی این پژوهش آن است که خودمختاری هوش مصنوعی نباید به معنای انتقال مسئولیت از انسان به ماشین تلقی شود؛ زیرا هوش مصنوعی فاقد شخصیت، اراده مستقل حقوقی و قابلیت تحمل مسئولیت به معنای سنتی است. بلکه خودمختاری باید به‌عنوان عاملی برای بازنگری در معیارهای انتساب مسئولیت مورد توجه قرار گیرد. به بیان دیگر، هرچه میزان استقلال و پیچیدگی سامانه افزایش می‌یابد، نظام حقوقی نیز باید از معیارهای سنتی مبتنی بر تقصیر فاصله گرفته و به سمت معیارهای عینی‌تر مانند مسئولیت مبتنی بر خطر، مسئولیت حرفه‌ای تشدیدشده و سازوکارهای جبرانی حرکت کند.

از منظر سیاست جنایی نیز، یافته‌های پژوهش نشان می‌دهد که مواجهه مؤثر با خسارات ناشی از هوش مصنوعی نیازمند عبور از سیاست جنایی واکنشی و حرکت به سوی سیاست جنایی پیشگیرانه است. در این رویکرد، هدف صرفاً مجازات یا جبران خسارت پس از وقوع حادثه نیست، بلکه کاهش احتمال وقوع زیان از طریق نظارت پیشینی، ارزیابی خطر، استانداردسازی فنی، شفافیت الگوریتمی و ایجاد نظام پاسخگویی فناورانه است.

بر این اساس، برای اصلاح نظام حقوقی ایران پیشنهادها زیر ارائه می‌شود:

۱. تدوین قانون جامع مسئولیت مدنی هوش مصنوعی

ضروری است قانون‌گذار ایرانی با الهام از تجربه اتحادیه اروپا، قانونی اختصاصی درباره مسئولیت ناشی از عملکرد سامانه‌های هوش مصنوعی تدوین کند. این قانون باید ضمن حفظ قابلیت اجرای قواعد عمومی مسئولیت مدنی، معیارهای ویژه‌ای برای سامانه‌های خودمختار، اشخاص مسئول، نحوه اثبات رابطه سببیت و حدود مسئولیت تعیین نماید.

در مواردی که تعیین مسئول دشوار است یا خسارت ناشی از نقص ناشناخته فناوری ایجاد شده است، صندوق جبران خسارت می تواند نقش حمایتی ایفا کند. منابع این صندوق می تواند از محل حق بیمه، عوارض فعالیت های پرخطر فناورانه یا مشارکت شرکت های فعال در حوزه هوش مصنوعی تأمین شود.

۶. توسعه الزامات شفافیت و قابلیت توضیح الگوریتم ها یکی از مهم ترین موانع رسیدگی قضایی به خسارات هوش مصنوعی، عدم امکان فهم فرآیند تصمیم گیری الگوریتمی است. بنابراین، تولیدکنندگان و ارائه دهندگان سامانه های هوشمند باید مکلف به ثبت داده های عملکرد، نگهداری سوابق تصمیم گیری و ارائه اطلاعات لازم برای بررسی قضایی باشند.

در نهایت، می توان گفت که آینده مسئولیت مدنی در عصر هوش مصنوعی نیازمند ایجاد تعادل میان دو ارزش بنیادین است: نخست، حمایت از حقوق و امنیت اشخاص در برابر خطرات فناوری؛ و دوم، حفظ امکان توسعه و بهره برداری مسئولانه از نوآوری های هوشمند. حقوق ایران برای دستیابی به این تعادل، باید از رویکرد صرفاً سنتی مسئولیت فاصله گرفته و با بهره گیری از تجربه نظام های پیشرو، چارچوبی ترکیبی از مسئولیت مدنی، سیاست جنایی پیشگیرانه و تنظیم گری فناورانه ایجاد کند.

۲. پذیرش نظام مسئولیت مبتنی بر خطر برای سامانه های پرخطر در حوزه هایی مانند پزشکی، حمل و نقل، امنیت و زیرساخت های حیاتی، که استفاده از هوش مصنوعی می تواند خسارات گسترده ایجاد کند، مناسب است مسئولیت مبتنی بر خطر جایگزین انحصاری نظریه تقصیر شود. در این صورت، شخصی که از فعالیت پرخطر هوش مصنوعی بهره اقتصادی می برد، مسئول کنترل و جبران پیامدهای آن خواهد بود.

۳. ایجاد نهاد ملی نظارت بر هوش مصنوعی تشکیل نهادی تخصصی برای ارزیابی، صدور مجوز، نظارت و کنترل سامانه های هوش مصنوعی می تواند نقش مهمی در سیاست جنایی پیشگیرانه ایفا کند. این نهاد باید اختیار بررسی میزان خطر سامانه ها، تعیین استانداردهای ایمنی و نظارت بر رعایت الزامات قانونی را داشته باشد.

۴. پیش بینی بیمه اجباری مسئولیت هوش مصنوعی با توجه به پیچیدگی و گستردگی خسارات احتمالی، بیمه مسئولیت می تواند ابزار مؤثری برای حمایت از زیان دیدگان باشد. در حوزه هایی که استفاده از هوش مصنوعی با خطر بالاتر همراه است، الزام به داشتن پوشش بیمه ای باید به عنوان شرط فعالیت قانونی پیش بینی شود.

۵. ایجاد صندوق جبران خسارت ناشی از هوش مصنوعی

