

## مسئولیت مدنی توسعه‌دهندگان و مدیران هوش مصنوعی با رویکرد کارآمدی اقتصادی

شقایق قنبریان<sup>۱</sup>، جواد حسین زاده<sup>۲</sup>، سید محمد آذین<sup>۳</sup>

### چکیده

هوش مصنوعی به عنوان پدیده‌ای نوظهور، به‌رغم مزایای قابل توجه، چالش‌های جدی در حوزه مسئولیت مدنی ایجاد کرده است. سیستم‌های هوش مصنوعی با ویژگی‌های پیچیدگی و استقلال در عملکرد، ماهیت سنتی رابطه سببیت و تقصیر را با تردید مواجه ساخته‌اند. از منظر عدالت اصلاحی، تدوین قواعد مسئولیت متناسب با ماهیت این سیستم‌ها امری ضروری است، چرا که رفتار مستقل آن‌ها پیش‌بینی‌پذیری را کاهش می‌دهد و دامنه خسارات احتمالی را گسترده‌تر می‌سازد؛ با این حال، حتی در فرض عملکرد کاملاً مستقل، نقش عامل انسانی در زنجیره توسعه، تولید و مدیریت قابل چشم‌پوشی نیست؛ از این رو، قواعد حاکم بر مسئولیت مدنی باید به‌گونه‌ای طراحی شود که انگیزه‌های مؤثری برای کاهش هزینه‌های اجتماعی ناشی از حوادث مرتبط با هوش مصنوعی ایجاد نمایند. مقاله حاضر با اتخاذ رویکرد کارآمدی اقتصادی، در پی پاسخ به این پرسش است که کدام نظریه مسئولیت برای هر یک از این دو گروه کارآمدتر است. یافته‌ها نشان می‌دهد مسئولیت محض برای توسعه‌دهندگان و نظریه تقصیر برای مدیران، به ترتیب می‌تواند زمینه‌ساز کارآمدی اقتصادی شود. این الگوی ترکیبی، ضمن ایجاد انگیزه‌های پیشگیرانه، توازن میان حمایت از زیان‌دیده و تشویق به نوآوری برقرار سازد.

**واژگان کلیدی:** هوش مصنوعی، مسئولیت مدنی، مسئولیت محض، نظریه تقصیر، کارآمدی اقتصادی، هزینه‌های اجتماعی، توسعه‌دهندگان، مدیران.

۱. دانشجوی دکتری رشته حقوق خصوصی، دانشکده حقوق، دانشگاه علم و فرهنگ: تهران، ایران  
shaghayeghghanbarian@yahoo.com

۲. دانشیار دانشکده حقوق، دانشگاه علم و فرهنگ، تهران، ایران (نویسنده مسوول)  
hoseinzadeh@usc.ac.ir

۳. مدیرگروه حقوق پزشکی، دانشگاه علم و فرهنگ، تهران، ایران  
azin@usc.ac.ir

## درآمد

هوش مصنوعی در کنار مزایای مختلفی که دارد، سبب ایجاد نگرانی در مورد خطرات ناشی از عملکرد آن شده است. یکی از سؤالات کلیدی در این زمینه این است که آیا نظام‌های حقوقی به اندازه کافی برای مقابله با زیان‌های مرتبط با هوش مصنوعی تکامل یافته‌اند یا خیر. در این باره گفته‌اند توجه به این نکته بسیار حائز اهمیت است که قوانین و مقررات حمایت از مصرف‌کنندگان در دوره‌ای تهیه شده است که تولیدکننده خطرات را کنترل می‌کند و بنابراین کاملاً مسئول زیان‌های ناشی از نقص محصول می‌باشد (رجبی، ۱۳۹۸: ۴۵۵). این ایده در دوره حاضر یعنی از نقطه‌نظر دوره پویاتر فناوری‌های جدید، منسوخ شده است، جایی که نه تنها تولیدکنندگان سنتی، بلکه توسعه‌دهندگان و مدیران هوش مصنوعی نیز ممکن است سبب زیان‌های ناشی از محصول باشند، هر چند که امکان کنترل خطرات فراهم نباشد. همین امر سبب توجه نظام‌های حقوقی به تدوین مقرراتی برای هوش مصنوعی گردیده و اتحادیه اروپا اخیراً در سیزدهم ژوئن سال ۲۰۲۴ دستورالعمل شماره (EU) ۱۶۸۹/۲۰۲۴<sup>۱</sup> را به منظور وضع قوانین هماهنگ در مورد هوش مصنوعی به تصویب رسانده است. منظور از توسعه دهنده، مطابق با بند سیزده دستورالعمل یاد شده هر شخص حقیقی یا حقوقی، از جمله یک مقام عمومی، سازمان یا ارگانی است که از یک سیستم هوش مصنوعی تحت اختیار خود استفاده می‌کند (به استثنای مواردی که سیستم هوش مصنوعی در جریان یک فعالیت شخصی غیر حرفه‌ای بسته به نوع سیستم هوش مصنوعی به کار می‌رود) و این استفاده ممکن است افراد دیگری غیر از توسعه دهنده را تحت تأثیر قرار دهد. در حوزه هوش مصنوعی مدیران نیز به دو دسته نزدیک (یا میانی)<sup>۲</sup> و مدیر پشتیبان قابل تقسیم

1. Regulation (EU) 2024/1689 of the European Parliament and of the Council. Regulation (EU) 2024/1689

قانون هوش مصنوعی اتحادیه اروپا اولین چارچوب جامع جهانی برای تنظیم هوش مصنوعی مبتنی بر رویکرد ریسک است. سیستمها به چهار سطح ریسک غیر قابل قبول (ممنوع)، بالا (با الزامات سختگیرانه)، محدود (نیازمند شفافیت) و حداقلی (بدون تعهد) طبقه بندی می شوند. این قانون شامل مانند تعهدات شفافیتی برای مدل های عمومی و جریمه های تا ۳۵ میلیون یورو یا هفت درصد گردش مالی جهانی است. نهادهای نظارتی جدید شامل دفتر هوش مصنوعی و هیئت هوش مصنوعی ایجاد شده اند.

۲. در بحث شرکت‌ها در تعریف مدیر میانی بیان شده مدیری است که تحت نظارت و هدایت مدیر

می‌باشند؛ مدیر پشتیبان که می‌تواند سبب مقدم در تأثیر نیز نامیده شود به صورت غیر مستقیم کنترل و هدایت عملکرد هوش مصنوعی را اعمال می‌کند و می‌تواند شامل توسعه دهنده نیز باشد اما مدیر نزدیک آخرین کنترل و نظارت را بر فعالیت هوش مصنوعی اعمال می‌کند.

نسبت به موضوع زیان‌های ناشی از هوش مصنوعی مسئولیت اشخاصی از جمله تولیدکننده، مالک، متصرف، توسعه دهنده، طراح و مدیر متصور می‌باشد؛ در این پژوهش، مسئولیت مدنی زیان‌های مرتبط با هوش مصنوعی و به ویژه تحولات اخیر اتحادیه اروپا در پرتو یک چارچوب حقوقی و اقتصادی با مرکزیت توسعه‌دهندگان و مدیران هوش مصنوعی مورد بحث قرار گرفته است. مزیت رویکرد حقوقی و اقتصادی این است که معیار روشنی (رفاه اجتماعی)<sup>۱</sup> را ارائه می‌کند که به انتخاب مبنای مناسب در زمینه مسئولیت مدنی کمک کرده و در عین حال مهم‌ترین اهداف چنین قوانینی یعنی بازدارندگی و تخصیص مسئولیت را در نظر می‌گیرد.

به‌طور مشخص سعی بر آن است که در این پژوهش به این سوال مهم پاسخ داده شود که از منظر نظریه‌های اقتصادی، تخصیص بهینه مسئولیت برای زیان‌های مرتبط با هوش مصنوعی به چه صورت است. این نوشتار ابتدا یک دیدگاه حقوقی و اقتصادی جهت تعیین قواعد بهینه برای بازیگران مختلف در طول زنجیره عمل هوش مصنوعی ارائه می‌دهد. سپس، بر اساس رویکردهای اقتصادی، به دنبال برجسته کردن مشکلات و ابهامات در مسئولیت مدنی ناشی از هوش مصنوعی و انتخاب نظریه احسن است. در این راستا، ابتدا موضوع با بحثی مختصر در مورد مفهوم و جایگاه هوش مصنوعی در نظام مسئولیت مدنی آغاز می‌گردد و سپس مدل اقتصادی مبنای

---

شرکت تجاری بوده و توسط وی منصوب می‌شود، اما کنترل و هدایت زیردستان خود را نیز بر عهده دارد و مسئول شمردن چنین مدیرانی در عمل و در رویه قضایی ایران نیز مشاهده شده است (بیگی حبیب‌آبادی و عیسایی تفرشی و شهنازی نیا، ۱۳۹۹).

۱. معیار رفاه اجتماعی از جمله نتیجه کارایی اقتصادی تلقی می‌شود، زیرا رفاه اجتماعی از طریق به‌کارگیری نظریه‌های ابزارگرا یعنی نظریه‌های اقتصادی، نظریه‌های مبتنی بر عدالت توزیعی و نظام‌های جایگزین مسئولیت مدنی محقق می‌شود؛ بر خلاف نظریه‌های مرسوم که از دیدگاه عدالت بین طرفین و حقوق خصوصی به مسئولیت مدنی نگاه می‌شود در نظریه‌های ابزارگرا مسئولیت مدنی ابزاری برای تحقق عدالت اجتماعی است و از دیدگاه حقوق عمومی بررسی می‌شود (بادینی، ۱۴۰۲: ۳۷۹). در واقع در نظریات ابزارگرا، مسئولیت مدنی به‌مثابه ابزاری است برای دستیابی به هدف‌هایی که از نظر اجتماعی مطلوب تلقی می‌شوند (رحیمی و طرف، ۱۳۹۷: ۷۳).

قانون مسئولیت مدنی را هنگام مواجهه با حوادث ارائه می‌دهد. در ادامه تجزیه و تحلیل خود را با توضیح این‌که چگونه این مدل اقتصادی پایه می‌تواند به تعیین مناسب‌ترین قواعد مسئولیت مدنی در خصوص زیان‌های مرتبط با هوش مصنوعی کمک کند، به پایان می‌رساند.

### ۱. مفهوم و جایگاه هوش مصنوعی در نظام مسئولیت مدنی

ویژگی استقلال در تصمیم‌گیری هوش مصنوعی سبب شده تا به عنوان یکی از مسئولیت‌های خاص از جانب حقوق‌دانان مورد بررسی قرار گیرد؛ در نتیجه، توجه و تعمق در نظام مسئولیتی متناسب با هوش مصنوعی و بدو تبیین مفهوم آن گریزناپذیر است.

#### ۱-۱. مفهوم هوش مصنوعی و چالش‌های بالقوه آن

همان‌طور که در بند دوم دامنه کاربرد توصیه‌نامه یونسکو<sup>۱</sup> بیان شده است، ارائه معنای واحد برای هوش مصنوعی به دلیل گسترده بودن قلمرو آن و پیشرفت مستمر تکنولوژی امکان‌پذیر نیست. با این حال، هوش مصنوعی به شکلی از فناوری اطلاق می‌شود که می‌تواند با عمل هوشمندانه و به شیوه‌ای «مشابه انسان» به اهداف معینی دست یابد. هوش مصنوعی با سیستم‌های از پیش برنامه‌ریزی شده و به دلیل توانایی آن در یادگیری، به صورت مستقل تصمیم می‌گیرد (Kingston, 2018: 2). مفهوم سیستم هوش مصنوعی نسبتاً گسترده است: اگر یک ماشین یا برنامه کامپیوتری بتواند با میزانی از استقلال عمل کند، صرف نظر از این‌که جزئی از محصول دیگر یا نرم‌افزاری مستقل باشد، می‌تواند سیستم هوش مصنوعی نامیده شود. در عین حال، عنصر استقلال کامل که توسط نظریه پردازان هوش مصنوعی تأیید شده است، بیانگر آن است که فرایندهای هوشمند می‌توانند بدون دخالت افراد دیگر دارای ادراک بوده و با فکر عمل کنند (Kārklīņš, 2020: 167).

از منظر کمیسیون اروپا هوش مصنوعی به سیستم‌هایی گفته می‌شود که

۱. توصیه‌نامه یونسکو درباره اخلاق هوش مصنوعی، مصوب ۲۳ نوامبر ۲۰۲۱، در بند دوم خود، ده اصل بنیادین شامل تناسب و عدم ورود ضرر، ایمنی و امنیت، انصاف و عدم تبعیض، پایداری، حق بر حریم خصوصی و حفاظت از داده‌ها، نظارت و تصمیم‌گیری انسانی، مسئولیت و پاسخ‌گویی، شفافیت و قابلیت توضیح دهی، آگاهی و سواد، و حاکمیت و همکاری چند ذی‌نفع و تطبیقی را به عنوان شالوده رویکرد حقوق بشری به اخلاق هوش مصنوعی معرفی می‌کند.

رفتار هوشمندانه را با تجزیه و تحلیل محیط خود و اقدام با درجه‌ای از استقلال برای دستیابی به اهداف خاص نشان می‌دهند. سیستم‌های فعال با هوش مصنوعی می‌توانند صرفاً نرم‌افزار باشند، که در دنیای مجازی عمل می‌کنند (به عنوان مثال دستیار صوتی، نرم‌افزار تجزیه و تحلیل تصویر، موتورهای جستجو، یا سیستم‌های تشخیص صدا و چهره)، اما هوش مصنوعی هم‌چنین می‌تواند در دستگاه‌های سخت‌افزاری (ربات‌های پیشرفته، خودروهای خودران، هواپیماهای بدون خلبان و ...) تعبیه شود؛ بنابراین بسیاری از فناوری‌های هوش مصنوعی برای کارآمدتر بودن به داده‌ها نیاز دارند. در نتیجه، می‌بینیم که پس از موفقیت عملکردی، این داده‌ها می‌توانند به بهبود و خودکارسازی روند تصمیم‌گیری در هوش مصنوعی کمک کنند (Yeung, 2019: 16)؛ بنابراین هوش مصنوعی واقعی اکنون به عنوان سیستمی که رفتار هوشمندانه را با تجزیه و تحلیل محیط خود و اقدام با سطحی از استقلال برای دستیابی به اهداف خاص نشان می‌دهد، تعریف شده است. هوش مصنوعی شامل ظرفیت یادگیری در حال توسعه مدل رفتاری خود (جانشینی دستورالعمل‌ها یا قوانین) از یافته‌ها و یا آمار گرفته شده از داده‌هایی است که به آن دسترسی دارد. به این معنی که هوش مصنوعی نوعی یادگیری عمیق مرتبط است که هم‌چنین می‌تواند از شبکه‌های عصبی مصنوعی بسیار بزرگ مدل شده در عملکرد سیستم عصبی انسان استفاده کند. ویژگی این هوش مصنوعی جدید در استقلال آن یا حتی آزادی کامل آن از طراح نهفته است (Céline Mangematin, 2020: 447). به همین دلیل همانند اشخاص حقیقی و حقوقی، هوش مصنوعی نیز ممکن است خطرهایی ایجاد کند که باعث تضرر دیگری شود و این ضرر ممکن است مادی یا غیر مادی باشد. به عنوان مثال، هوش مصنوعی در دستگاه‌های پزشکی یا وسایل نقلیه خودران ممکن است به افراد یا دارایی‌ها خسارت فیزیکی وارد کند و باعث درد یا رنج دیگری شود. سیستم‌های هوش مصنوعی هم‌چنین می‌توانند زیان‌های اقتصادی متفاوتی ایجاد کنند. برای مثال، نقص در یک سیستم مدیریت جاده ممکن است منجر به ازدحام شده که این امر می‌تواند به زیان‌های متعدد، از قبیل فوت فرصت انعقاد یک قرارداد تجاری منتج گردد.

ظرفیت هوش مصنوعی برای زیان رساندن نیاز به بررسی فوری راه‌های

کاهش چنین حوادثی دارد. اولین مساله در این خصوص آن است که آیا انسان مسؤول تلقی شود؟ مفسران این سوال را مطرح کرده‌اند که آیا سیستم‌های هوش مصنوعی باید دارای شخصیت حقوقی باشند تا آن‌ها را در قبال زیان‌های مرتبط مسؤول بدانند. با این حال، این رویکرد نه تنها به این دلیل که نظارت انسانی را نادیده می‌گیرد، بلکه به دلیل ادعای غیر ممکن بودن آن، زمانی که مثلاً منابع مالی مورد نیاز برای جبران خسارت را در نظر می‌گیریم، مورد انتقاد قرار گرفته است (ذاکری‌نیا، ۱۴۰۲: ۱۴۱). برخی گفته‌اند که انتساب شخصیت مصنوعی و «مسؤولیت شخصی» به سیستم‌های هوش مصنوعی فاقد مبنا و منشأ قانونی است (نظریور و ایزدی‌فرد و محسنی دهکلانی، ۱۴۰۰: ۱۸۵). افزون بر این، به‌طور کلی توافق شده است که نظارت انسانی ضروری است (ولی‌پور و اسماعیلی، ۱۴۰۰: ۱۱) و در این مرحله، سیستم‌های هوش مصنوعی نباید بازیگران مستقلى از نقطه‌نظر قانونی باشند. نقطه شروع این مقررات باید به انسان‌ها و نهادهایی که برنامه‌های هوش مصنوعی را تولید و کنترل می‌کنند مربوط باشد. صرف‌نظر از پذیرش یا رد نظریه شخصیت حقوقی (یا مصنوعی) که مستقلاً نیازمند بررسی می‌باشد و به نظر ما در این بین پذیرش شخصیت حقوقی برای هوش مصنوعی قابل دفاع است؛ با این حال باز هم ضروری است که برای بازیگران انسانی درگیر در این حوزه، انگیزه‌هایی برای رفتار مناسب فراهم شود. البته برخی از ویژگی‌های هوش مصنوعی، تخصیص مسؤولیت را در این حوزه به یک کار چالش برانگیز تبدیل می‌کند.

اول آن‌که، با عنایت به دخالت بازیگران مختلف مانند ذی‌نفعان و تولیدکنندگان متعدد، تشخیص عاملی که باعث ورود زیان شده است به یک کار پیچیده تبدیل گردیده است (گندم‌کار و صالحی مازندرانی و حمیدی، ۱۴۰۰: ۲۴۳). اشخاص متعددی در ایجاد یک سیستم هوش مصنوعی مستقل یا محصولی که چنین سیستمی را تعبیه می‌کند، مشارکت می‌کنند، اما در این میان اشخاصی هم هستند که بر عملکرد یک سیستم هوش مصنوعی کنترل دارند. تصمیم‌گیری درباره این‌که چگونه رفتار هر یک از این بازیگران باید کنترل و از نظر حقوقی توصیف شود امر دشواری می‌نماید.

دوم، فرایند تصمیم‌گیری یک سیستم هوش مصنوعی می‌تواند مبهم باشد. بر

خلاف ماشین‌های قابل برنامه‌ریزی، که عملکرد آن‌ها در محدودیت‌های نرم‌افزار قابل پیش‌بینی است، یک سیستم هوش مصنوعی پیشرفته، تصمیم‌گیرنده‌ای غیر شفاف است (Kärklīņš, 2020: 168). این تصمیم‌گیری غیر شفاف ممکن است خطر بیشتری ایجاد کند که برای بازیگران حوزه هوش مصنوعی ناشناخته است و در نتیجه ارزیابی احتمال وقوع آسیب را پیچیده سازد (Arcila, 2024: 3). علاوه بر این، تعداد موقعیت‌هایی که در آن هر یک از طرفین رفتار مناسبی داشته‌اند اما با این وجود حادثه رخ می‌دهد، می‌تواند افزایش یابد.

سوم، عملکرد یک سیستم هوش مصنوعی پیشرفته به‌طور کامل یا تا حدی مستقل است (Vladeck, 2014: 123)؛ این استقلال در عمل رابطه بین انسان و ماشین و امکان انتساب عمل ماشین به انسان را غیر شفاف می‌سازد. وضعیت مزبور هر چند سبب طرح نظریه شخصیت حقوقی هوش مصنوعی شده است، اما با ادعای عدم انطباق کامل ممیزات شخص حقیقی و حقوقی بر هوش مصنوعی (ذاکری‌نیا، ۱۴۰۲: ۱۴۰) و وجود دشواری‌های جبران خسارت توسط هوش مصنوعی مورد انتقاد قرار گرفته است.

## ۲-۱. مسئولیت غیر قراردادی

همان‌طور که بیان شد سیستم‌های مختلف هوش مصنوعی با فعالیت‌های متعددی که می‌توانند داشته باشند ممکن است سبب بروز خسارت گردند و دلیل این امر می‌تواند ناشی از عدم دقت در برنامه‌ریزی، نقص درونی سخت‌افزاری یا حتی بنا به دلایل غیر قابل کشف باشد (Bartlett, 2023: 241) که در فرض استقلال این سامانه‌ها خسارت وارده قابل پیش‌بینی نمی‌باشد تا کاربر بتواند از بروز آن جلوگیری نماید (گندم‌کار و صالحی مازندرانی و حمیدی، ۱۳۹۹: ۲۴۳). کاهش ریسک و جبران ضررهای وارده لزوماً به ابزار قانونی نیاز ندارد و در صورتی که طرفین بتوانند بدون چالش با یکدیگر مذاکره نموده و همه آن‌ها اطلاعات کامل و لازم برای شناسایی اسباب مختلف را در اختیار داشته باشند، می‌توان ریسک و ضرر مرتبط را به‌طور بهینه از طریق چانه‌زنی بین بازیگران مختلف تخصیص داد. در شرایط ایده‌آل، هرگونه مداخله قانونی غیر ضروری و ناکارآمد است، زیرا می‌تواند منجر به تحمیل هزینه‌های اضافی شود (بادینی، ۱۴۰۲: ۴۲۲)؛ با این حال ابزارهای حقوقی زمانی ضروری

می‌شوند که انعقاد قرارداد پیش از ورود ضرر نتواند به یک نتیجه مطلوب در توزیع خسارت و تخصیص مسئولیت منجر شود، اگر هزینه‌های تخصیص مسئولیت بالا باشد یا عدم تقارن اطلاعاتی وجود داشته باشد، ممکن است طرفین نتوانند در مورد تخصیص کارآمد مسئولیت به توافق برسند. هم‌چنین، خطرهای ناشی از کالاهای صنعتی ممکن است از بهای داد و ستد آن فراتر باشد و این امر در مورد سیستم‌های هوش مصنوعی نیز اتفاق می‌افتد. بازیگران مختلف در طول زنجیره عمل هوش مصنوعی ممکن است اشخاص حقوقی متفاوتی را نمایندگی کرده و دارای مقادیر متفاوتی از اطلاعات در مورد هوش مصنوعی باشند. علاوه بر این، در برخی موارد زیان‌دیدگان اشخاص ثالثی هستند که با هیچ‌یک از بازیگران هوش مصنوعی روابط قراردادی ندارند اما متضرر گردیده‌اند و لذا مسئولیت را باید از قرارداد خارج کرد (رجبی، ۱۳۹۸: ۴۵۴)؛ هوش مصنوعی شامل بازیگران و عناصر مختلفی می‌شود که در طراحی سیستم در طول چرخه عمر آن شرکت می‌کنند، آن را برنامه‌ریزی می‌کنند، تصمیم می‌گیرند که چه زمانی و برای چه چیزی استفاده شود و بر آن نظارت دارند. مطابق با مواد ۱۶ به بعد دستورالعمل شماره مقرر اتحادیه اروپا، درباره وضع قوانین هماهنگ برای هوش مصنوعی (قانون هوش مصنوعی اروپا) اتحادیه اروپا بازیگران حوزه هوش مصنوعی که برای آن‌ها تعهداتی پیش‌بینی شده است عبارتند از: ارائه‌دهندگان، واردکنندگان، توزیع‌کنندگان، توسعه‌دهندگان و نهادهای مطلع.<sup>۱</sup> شایان ذکر است دستورالعمل مذکور بیشتر از هوش مصنوعی پرخطر یاد نموده است و از مفاد آن می‌توان پی برد که درصدد تعیین تعهداتی برای برخی از دست‌اندرکاران حوزه تولید و به‌کارگیری هوش مصنوعی بوده تا تعیین بازیگران و طرح مسئولیت آن‌ها. ماده ۲۶ دستورالعمل یاد شده مقرر می‌دارد: «برای معرفی مجموعه‌ای متناسب و مؤثر از قوانین الزام‌آور برای سیستم‌های هوش مصنوعی، باید از یک رویکرد مبتنی بر ریسک که به روشنی تعریف شده باشد، پیروی شود. این رویکرد باید نوع و محتوای چنین قوانینی را با شدت و دامنه خطراتی که سیستم‌های هوش مصنوعی می‌توانند ایجاد کنند، متناسب کند؛ بنابراین، لازم است برخی از شیوه‌های غیر قابل قبول هوش مصنوعی ممنوع شود، الزاماتی برای سیستم‌های هوش مصنوعی پرخطر و تعهداتی

---

1. Notified Body

برای اشخاص مربوطه وضع شود و تعهدات شفافیت برای برخی از سیستم‌های هوش مصنوعی پیش‌بینی گردد». هم‌چنین، شق اول ماده ۷ دستورالعمل مزبور بیان داشته است: «به منظور تضمین سطح بالایی از حفاظت از منافع عمومی در رابطه با سلامت، ایمنی و حقوق اساسی، باید قوانین مشترکی برای سیستم‌های هوش مصنوعی پرخطر وضع شود ...».

دخالت افراد یا بازیگران متعدد در توسعه، استقرار و بهره‌برداری از سیستم‌های هوش مصنوعی تخصیص مسئولیت و پاسخ‌گویی برای نتایج هوش مصنوعی را پیچیده می‌کند (Arcila, 2024: 3)؛ از این‌رو مسئولیت مدنی غیر قراردادی به‌ویژه زمانی که نمی‌توان قبل از حادثه، فرایند مذاکره را ترتیب داد ابزاری حیاتی در ارائه انگیزه‌هایی به دست‌اندرکاران حوزه تولید و به‌کارگیری هوش مصنوعی برای رفتار مناسب و تبیین نظام مسئولیتی متناسب با هوش مصنوعی به‌ویژه با در نظر گرفتن مدل‌های اقتصادی مسئولیت است.

## ۲. مدل اقتصادی مسئولیت و تعیین قواعد مرتبط با آن در حوزه هوش

### مصنوعی

به طور کلی، مقررات حاکم بر مسئولیت مدنی، مبتنی بر نظریه تقصیر یا نظریه خطر است (صفایی، ۱۴۰۰: ۶۱؛ بادینی، ۱۴۰۲: ۷۸). حالت کارایی اقتصادی که حاصل نگرش اقتصادی به مسئولیت مدنی است، بیان می‌دارد که تنها کاربرد نظام مسئولیت مدنی این است که به فعالیت‌های خطرآفرین، تنها زمانی اجازه فعالیت داده شود که ارزش اجتماعی این گونه فعالیت‌ها به اندازه‌ای باشد که خطرات آن را توجیه کند (بادینی، ۱۴۰۲: ۳۸۸). بر اساس تئوری حقوق مسئولیت مدنی و اقتصاد، رفاه اجتماعی به عنوان معیاری برای تعیین این‌که کدام نوع قاعده مسئولیت برای برخورد با حوادث مناسب است عمل می‌کند. رفاه اجتماعی نه تنها تحت تأثیر اثر بازدارنده قواعد مسئولیت مدنی است، بلکه به اثر تخصیص مسئولیت نیز بستگی دارد به شکلی که، استفاده منطقی از یک سازکار توزیع عادلانه ریسک سبب تشویق نوآوری و کاهش هزینه عدم اطمینان برای همه فعالان اقتصادی و افزایش اعتماد مصرف‌کننده می‌شود (ذاکری‌نیا، ۱۴۰۲: ۱۳۷). در این بخش، نحوه شناسایی قاعده مطلوب با در نظر گرفتن این دو عامل توضیح داده شده است.

## ۲-۱. بازدارندگی به عنوان عاملی در تعیین قواعد کارآمد مسؤولیت

### مدنی

از منظر اقتصادی هدف اولیه قواعد مسؤولیت، بازدارندگی است (بادینی، ۱۴۰۲: ۳۵۷). در نظریه‌های مبتنی بر بازدارندگی، هدف مسؤولیت مدنی این است که از طریق انعکاس هزینه‌های حوادث در قیمت فعالیت‌ها، افراد را وادار نماید تا مجموع هزینه‌های حوادث و هزینه‌های پیش‌گیری از حوادث را به حداقل برسانند (سلیمی و پارساپور، ۱۴۰۰: ۵۷). در یک حادثه دو جانبه، که هر دو عامل زیان و زیان‌دیده در وقوع حادثه سهیم هستند، در صورتی که یک قاعده مسؤولیت بتواند رفتار همه افراد را کارآمد کند، می‌توان به یک نتیجه اجتماعی به صرفه برای طرفین دست یافت. این که آیا یک قاعده حقوقی می‌تواند به نتیجه اجتماعی کارآمد منتج شود به دو متغیر بستگی دارد.

از یک سو در مواردی که احتیاط از جانب هر دو طرف ممکن است، هم مسؤولیت محض با امکان استناد به تقصیر مشارکتی زیان‌دیده و هم مسؤولیت مبتنی بر تقصیر با یا بدون امکان استناد به تقصیر مشارکتی زیان‌دیده با و داشتن طرفین به در پیش گرفتن احتیاط و مراقبت لازم، باعث به وجود آمدن کارایی اقتصادی می‌شود (بادینی، ۱۴۰۲: ۴۴۵). رفاه اجتماعی متأثر از اقدامات پیش‌گیرانه طرفین است. یک قاعده مسؤولیت مطلوب به طرفین انگیزه می‌دهد تا سطح بهینه مراقبت را انجام دهند. این نقطه‌ای است که هزینه‌های نهایی یک اقدام احتیاطی با کاهش نهایی زیان مورد انتظار برابری می‌کند. با تنظیم سطح مراقبت لازم در این سطح بهینه، یا مسؤولیت محض (با دفاع از تقصیر مشارکتی) یا یک قاعده مبتنی بر تقصیر می‌تواند همه طرفین را وادار به انجام اقدامات احتیاطی بهینه کند.

از سوی دیگر زمانی که انتظار می‌رود همه طرفین سطح بهینه مراقبت را در هر یک از گزینه‌های قاعده مسؤولیت اتخاذ کنند، کنترل سطح فعالیت ضروری می‌شود (حکمت‌نیا و محمدی و واثقی، ۱۳۹۸: ۲۴۶). در صورتی که دست‌اندرکاران حوزه تولید و به‌کارگیری هوش مصنوعی سطح فعالیت خود را کنترل کنند، رفاه اجتماعی می‌تواند بیشتر بهبود یابد (بادینی، ۱۴۰۲: ۴۴۶). عواملی که بر سطح فعالیت تأثیر دارند (مانند سود شخصی و شدت و فراوانی یک فعالیت) به ندرت توسط افراد

خارجی قابل مشاهده است تا سطح فعالیت خود را بهینه کنند. در حالت ایده‌آل، قواعد مسئولیت همه اشخاص ذی‌ربط را تشویق می‌کند تا سطح فعالیت بهینه را اتخاذ کنند. با این حال، مدل‌های اقتصادی ثابت کرده‌اند که هیچ‌یک از قواعد موجود نمی‌تواند هر دو عامل زیان و زیان‌دیده را به بهینه‌سازی سطح فعالیت خود به‌طور هم‌زمان وادار سازد. علت این امر آن است که مسئولیت، زمانی که همه طرف‌ها کوشا بوده‌اند، باید بر مبنای نظریه خطر به هر یک از متخلفان یا بر مبنای نظریه تقصیر به زیان‌دیدگان تخصیص داده شود.

در این راستا، مسئولیت محض و قواعد مبتنی بر تقصیر منجر به نتایج متفاوتی در بهینه‌سازی سطح فعالیت می‌شود. از آن‌جا که طرفین باید مسئولیت مدنی را با توجه به این‌که همه طرف‌ها آمادگی دارند، درونی کنند، با اعمال مسئولیت محض برای متخلفان خاص، انگیزه‌ای برای بهینه‌سازی بیشتر سطح فعالیت خود برای به حداقل رساندن مسئولیت تحمیل شده توسط آن‌ها خواهد داشت. در مقابل، اگر متخلفان خاص مشمول مسئولیت مبتنی بر تقصیر باشند، اگر سطح مراقبت لازم را اتخاذ کرده باشند، مسئولیتی ندارند. از این‌رو، تحت قاعده مبتنی بر تقصیر، متخلفان تمایل به اتخاذ سطح بالای فعالیت برای به حداقل رساندن فایده شخصی خود دارند.

در جایی که هوش مصنوعی ممکن است خطرناک باشد اعمال مسئولیت بر توسعه‌دهندگان آن سبب بازدارندگی است؛ زیرا توسعه‌دهندگان هستند که از این معاملات نفع می‌برند (رجبی، ۱۳۹۸: ۴۵۶)؛ بنابراین از منظر بازدارندگی، ارزش مسئولیت محض این است که برای یک طرف انگیزه اضافی جهت کنترل فعالیت‌های خود فراهم کند، به ویژه زمانی که چنین فعالیت‌هایی دارای درجه بالایی از خطر هستند و ارزش افزوده کمی دارند. این دیدگاه یعنی پیش‌بینی مسئولیت محض با رویکرد علم احتمالات نیز قابل سنجش است و از آن‌جا که احتمال گریز از جبران خسارت اصولاً متصور نیست می‌تواند کارایی اقتصادی را در پی داشته باشد (کریمیان، ۱۴۰۱: ۱۳). وانگهی آن‌جا که سود کنترل سطح فعالیت، حاشیه‌ای است، به منظور جلوگیری از محدود کردن فعالیت‌های ارزشمند، قواعد مبتنی بر تقصیر ترجیح داده می‌شود. شایان ذکر است، مطابق با بند چهل و چهار دستورالعمل شماره

۱۶۸۹ مصوب ۲۰۲۴ اتحادیه اروپا: «از آنجا که نگرانی‌های جدی در مورد مبنای علمی سیستم‌های هوش مصنوعی با هدف شناسایی یا استنتاج احساسات وجود دارد، به‌ویژه از آنجا که بیان احساسات در فرهنگ‌ها و موقعیت‌ها و حتی در یک فرد به طور قابل توجهی متفاوت است و چنین سیستم‌هایی می‌تواند منجر به برخورد زیان‌بار یا نامطلوب با افراد خاص یا کل گروه‌های آن‌ها شود؛ بنابراین عرضه در بازار، راه‌اندازی یا استفاده از سیستم‌های هوش مصنوعی که برای تشخیص وضعیت عاطفی افراد در موقعیت‌های مربوط به محل کار و تحصیل استفاده می‌شوند، باید ممنوع شود. این ممنوعیت نباید سیستم‌های هوش مصنوعی را که صرفاً به دلایل پزشکی یا ایمنی در بازار عرضه می‌شوند، مانند سیستم‌هایی که برای استفاده درمانی در نظر گرفته شده‌اند، پوشش دهد»<sup>۱</sup>. در چنین مواردی موضوع استقلال هوش مصنوعی مطرح نمی‌شود و خسارات وارده عرفاً متناسب به عامل زیان است که اصولاً توسعه دهنده می‌باشد.

## ۲-۲. تخصیص مسئولیت به عنوان عاملی در تعیین قواعد کارآمد مسئولیت

### مدنی

از منظر اقتصادی، کارکرد یک قاعده کارآمد نه تنها متأثر از اثر بازدارنده آن است، بلکه تحت‌تأثیر تخصیص مسئولیت نیز قرار می‌گیرد. از دیدگاه عده زیادی از

1. Recital 44 of Regulation (EU) 2024/1689:(44) There are serious concerns about the scientific basis of AI systems aiming to identify or infer emotions, particularly as expression of emotions vary considerably across cultures and situations, and even within a single individual. Among the key shortcomings of such systems are the limited reliability, the lack of specificity and the limited generalisability. Therefore, AI systems identifying or inferring emotions or intentions of natural persons on the basis of their biometric data may lead to discriminatory outcomes and can be intrusive to the rights and freedoms of the concerned persons. Considering the imbalance of power in the context of work or education, combined with the intrusive nature of these systems, such systems could lead to detrimental or unfavourable treatment of certain natural persons or whole groups thereof. Therefore, the placing on the market, the putting into service, or the use of AI systems intended to be used to detect the emotional state of individuals in situations related to the workplace and education should be prohibited. That prohibition should not cover AI systems placed on the market strictly for medical or safety reasons, such as systems intended for therapeutic use.

طرفداران نگرش اقتصادی به حقوق، مهم‌ترین هدف درونی کردن هزینه‌های خارجی حوادث، کارایی تخصیصی بوده که امری غیر از هدف بازدارندگی اقتصادی می‌باشد. مسئولیت مدنی از طریق درونی کردن هزینه‌های اجتماعی حوادث باعث می‌شود تا کالاها و خدمات به گونه‌ای که از نظر اجتماعی مطلوب و کارا است، ارائه شود (بادینی، ۱۴۰۲: ۴۱۴). پس از وقوع حادثه، رفاه اجتماعی در صورتی افزایش می‌یابد که مسئولیت تخصیص یابد. از این رو، مقنن باید ترجیحات اشخاص مختلف و دسترسی احتمالی به سازکارهای تخصیص مسئولیت را به دقت ارزیابی کند. از منظر اقتصادی، فلسفه وجودی مسئولیت مدنی تحقق سازکار درونی کردن هزینه‌های اجتماعی حوادث از طریق تحمیل مسئولیت بر مسبب ورود زیان است زیرا هدف درونی کردن هزینه‌های خارجی، کارایی تخصیص است و گرنه هزینه تولید در قیمت‌گذاری در نظر گرفته شده و توسعه هوش مصنوعی به دلیل ارزان بودن، بیش از آن میزانی که از لحاظ اجتماعی مطلوب است، افزایش یافته و در نتیجه سطح کلی رفاه اجتماعی کاهش می‌یابد (بادینی، ۱۴۰۲: ۴۱۵ - ۴۱۶).

هم‌چنین عوامل دیگری نیز وجود دارند که در تخصیص مسئولیت نقش دارند. این عوامل شامل بیمه و قراردادهای مشارکت در مسئولیت است (Vladeck, 2014: 130). بند یک ماده ۳۱ دستورالعمل شماره ۱۶۸۹ مصوب ۲۰۲۴ اتحادیه اروپا کشورهای عضو را موظف به تأسیس یک شخصیت حقوقی تحت عنوان نهاد مطلع نموده است که این سازمان از جمله مرجع تأیید و صدور گواهی برای هوش مصنوعی ارائه شده به عنوان یک کالا به صورت موردی می‌باشد. بند نه همان ماده بیان داشته: «نهادهای مطلع باید بیمه مسئولیت مناسب را برای فعالیتهای خود با موضوع ارزیابی انطباق هوش مصنوعی با مقررات منعقد کنند، مگر این که کشور عضو که نهاد مزبور در آن تأسیس شده است مسئولیت را بر عهده گیرد یا آن کشور عضو مستقیماً مسؤول ارزیابی انطباق باشد»<sup>۱</sup>. از این امر می‌توان پی برد که در فروض تخصیص مسئولیت و بازدارندگی باید قائل به طرح مسئولیت اشخاص باشیم؛ زیرا مشخص نیست که قواعد عمومی در خصوص مسئولیت ماشین‌ها امکان اعمال داشته

1. A notified body shall be established under the national law of a Member State and shall have legal personality.

باشد (رجبی، ۱۳۹۸: ۶۲) و این رویکرد در هر حال از نظر مصالح اجتماعی نیز مستحسن است، زیرا مسؤول شناختن محصول هوشمند بر فرض استقلال عملکرد نیز سبب نفی مسؤولیت اشخاص نمی‌شود.

### ۳. نقش مدل‌های اقتصادی در تعیین نظام مسؤولیت هوش مصنوعی

در گذشته خطرات مربوط به محصولات ناشی از عمل تولیدکنندگان و مصرف‌کنندگان تلقی می‌شد. از یک سو، خطر می‌تواند در فرآیند ساخت محصول ایجاد شود و از سوی دیگر، رفتار مصرف‌کنندگان و بی‌دقتی آن‌ها می‌تواند منجر به آسیب شود. مصرف‌کنندگان و تولیدکنندگان در هر صورت می‌توانستند به راحتی یکدیگر را در بازار ملاقات کنند و مصرف‌کنندگان می‌توانستند کیفیت یک محصول را به صورت بصری ارزیابی کنند؛ بنابراین، دو طرف عملاً می‌توانستند مسؤولیت را با توافق یکدیگر تخصیص دهند. با وقوع انقلاب صنعتی وضعیت تغییر نمود و در حال حاضر نه تنها تولیدکنندگان و مصرف‌کنندگان اصلی، بلکه سایر بازیگران اعم از واردکنندگان، عرضه‌کنندگان، توسعه‌دهندگان و مدیران نیز می‌توانند در وقوع حوادث نقش داشته باشند. در نتیجه تخصیص مسؤولیت با استفاده از نظام مسؤولیت قراردادی، پیچیده و پرهزینه است. این پیش‌زمینه‌ای برای معرفی نظام‌های مسؤولیت غیر قراردادی خاص در رابطه با زیان‌های مربوط به تولید محصولات بود. تمام اشخاص مرتبط، مسؤول به‌کارگیری متصدیان واجد شرایط برای آزمایش ایمنی محصولات، بررسی قابلیت اطمینان فنی آن‌ها و گزارش‌های ریسک در فرآیند تولید هستند. با این حال، دادگاه‌ها، قانون‌گذاران و مصرف‌کنندگان نمی‌توانند به چنین اطلاعاتی به‌طور یکسان دسترسی داشته باشند. با توجه به خطر جدی ناشی از فعالیت تولیدکنندگان، ایجاد انگیزه برای آن‌ها به منظور بهینه‌سازی سطح فعالیت‌شان بسیار مهم است. لازم به ذکر است در نگاه شرکت‌های بین‌المللی نیز صرف تولید کافی نبوده و باید کالای سرمایه‌ای متناسب با دانش فنی تولید شده و در نهایت با ترکیب این دو به همراه مهارت‌های فنی، محصول نهایی را وارد بازار نمایند (السان، ۱۳۸۵: ۲۰۰). از آنچه گفته شد می‌توان پی برد که تنظیم معاهدات در بعد بین‌المللی و تدوین مقررات در حوزه داخلی جهت وضع برخی تعهدات برای دست‌اندرکاران حوزه هوش مصنوعی در زمینه انتقال، صادرات و فروش تکنولوژی به منظور حفظ بازار و

توسعه استفاده از هوش مصنوعی ضروری است.

افزون بر این، در مقایسه با مصرف‌کنندگان و سایر اشخاص، تولیدکنندگان توانایی بیشتری برای تحمل مسئولیت دارند. هم‌چنین، تولیدکنندگان با تحمل مسئولیت، بدون اشاره به سازکار تغییر مسئولیت مدنی، می‌توانند زیان‌های بالقوه را از طریق قیمت‌گذاری بین مصرف‌کنندگان توزیع کنند (سلیمی و پارساپور، ۱۴۰۰: ۴۸). مسئول شناختن تولیدکنندگان باید منتج به کاهش هزینه‌های حوادث و بالا رفتن سطح رفاه اجتماعی شود.

امروزه از آن‌جا که اشخاص متعددی در زنجیره فعالیت هوش مصنوعی درگیر هستند، وضعیت توزیع‌کنندگان بعدی و اشخاص غیر مصرف‌کننده نیز باید در نظر گرفته شود. اشخاص اخیر مسئول کاهش هزینه‌های حوادث با کنترل رفتار خود هستند؛ بنابراین، قواعد مسئولیت نیز باید انگیزه‌هایی را برای این اشخاص در طول زنجیره تامین فراهم کند. تولیدکنندگان اصلی می‌توانند ویژگی‌های اساسی محصولات را تعریف کنند، اشخاص مرتبط دیگر قادر به این امر نیستند. بیشتر خطرات ناشی از تقصیر آن‌ها در حمل و نقل یا استفاده از محصول است و به این دلایل، قواعد مبتنی بر تقصیر معمولاً برای این اشخاص اعمال می‌شود؛ بنابراین، در این موارد که سایر شرایط نیز در حالت معمول قرار دارد، تقصیر می‌تواند مبنایی برای انتساب زیان‌های وارده به مدیر یا مصرف‌کننده باشد. البته بدیهی است در چنین مواردی مصرف‌کننده در صورت وجود شرایط می‌تواند به توسعه‌دهندگان که دارای مسئولیت محض می‌باشند، مراجعه نماید. هم‌چنین، مسئولیت محض ممکن است با الزامات بیمه اجباری همراه باشد که خطر عدم پرداخت خسارت را برای زیان‌دیده کاهش می‌دهد.

همان‌طور که بیان گردید در بسیاری از موارد چندین شخص درگیر ارائه یک محصول یا خدمات هستند و اصولاً وجود خسارت در دادگاه به راحتی قابل اثبات است، اما این سوال باقی می‌ماند که چه کسی باید خسارت وارده را جبران کند و تا چه حد مسئولیت دارد. در این راستا، برخی بیان داشته‌اند که با توجه به استثنایی بودن مسئولیت تضامنی، مسئولیت اسباب متعدد اشتراکی است (ذاکری‌نیا، ۱۴۰۲: ۱۴۲). در این مورد می‌توان دو حالت را تصور نمود. یک حالت آن است که نظارت

و مراقبت جمعی است؛ یعنی ارائه مراقبت توسط یک طرف جایگزین کاملی برای مراقبت ارائه شده توسط شخص دیگر است. در حالت دوم، مراقبت از سوی همه طرفین ضروری است؛ یعنی ارائه مراقبت توسط یکی از اشخاص ذی‌مدخل مکملی برای مراقبت ارائه شده توسط طرف دیگر است. در چنین مواردی تشخیص اصل و میزان تقصیر این اشخاص معلوم نیست؛ بنابراین، می‌توان فرض کرد که هر دو به طور متقارن در خطر سهیم هستند و لذا مدعی می‌تواند هم‌زمان برای تمام خسارت به هر دو مراجعه کند اما بدیهی است از هر دو تنها سقف خسارت قابل مطالبه دریافت می‌شود.

#### ۴. مسؤولیت اشخاص مرتبط

اکنون ضرورت دارد نتایج نظری مطالب بیان شده در بندهای پیشین در زمینه هوش مصنوعی اعمال شود. در یک زنجیره فعالیت مرتبط با یک سیستم هوش مصنوعی، چندین طرف می‌توانند با مشارکت در توسعه یک سیستم هوش مصنوعی یا با بهره‌برداری از چنین سیستمی خطر ایجاد کنند. توانایی اقدام نسبت به پیش‌گیری و ظرفیت تحمل مسؤولیت در مورد توسعه‌دهندگان و سپس مدیران مورد بررسی قرار می‌گیرد.

#### ۴-۱. مسؤولیت توسعه‌دهندگان سیستم‌های هوش مصنوعی

به طور سنتی تصور می‌شد که تنها تولیدکنندگان اشیای فیزیکی ممکن است خطراتی را ایجاد کنند. با این حال، این امر دیگر در عصر دیجیتال صادق نیست، زیرا نرم‌افزار مستقل، که در هیچ محصولی تعبیه نشده است، ممکن است خطراتی را نیز به همراه داشته باشد؛ لذا، اشاره به تولیدکنندگان اشیای فیزیکی تا حدودی منسوخ شده است. با این حال، موضوع این نیست که آیا ما باید توسعه‌دهندگان سیستم‌های هوش مصنوعی را «تولیدکننده» بنامیم بلکه سوال نظری اساسی‌تر و مهم‌تر این است که آیا مسؤولیت محض یا قواعد مبتنی بر تقصیر باید برای همه توسعه‌دهندگان اعمال شود؟ برای پرداختن به این موضوع، باید به معیارهایی که در بند پیش توضیح داده شد برگردیم.

در جایی که دخالت انسان در تصمیم‌گیری ماشین بسیار مشهود است، نیازی به بررسی مجدد قواعد مسؤولیت مدنی نیست. هر شخصی که در توسعه

ماشین نقش داشته باشد و به طراحی تصمیم‌گیری آن کمک کند، عرفاً مسئول اعمال غیر قانونی اعم از عمدی و غیر عمدی است که توسط ماشین انجام می‌شود (Kingston, 2018: 7). دلیل این امر آن است که سیستم ماشینی، علی‌رغم پیچیدگی، فاقد شخصیت حقوقی می‌باشد و آن‌ها نمایندگان یا ابزارهای سایر نهادهایی هستند که به عنوان افراد، شرکت‌ها، یا سایر «اشخاص» حقوقی ممکن است طبق قانون در قبال اعمال خود پاسخ‌گو باشند (Vladeck, 2014: 121). البته در خصوص مواردی که به شکل مستقل عمل می‌نمایند هنوز چالش‌ها رفع نشده و لذا مباحث مطرح در این نوشتار چنین مصادیقی را نیز در برمی‌گیرد. در عصر هوش مصنوعی، ویژگی‌های دیجیتالی این توانایی را دارند که بر کارکرد و عملکرد یک محصول تأثیر بگذارند. نقش یک سیستم هوش مصنوعی حتی ممکن است مهم‌تر از دستگاه فیزیکی باشد؛ زیرا دارای اطلاعاتی است که به راحتی توسط افراد خارجی قابل دسترسی یا فهم نیست. به‌ویژه، ابزار به دست آمده از یک ساختار الگوریتمی و ریسک مربوط به آن ممکن است تنها توسط خود توسعه‌دهنده به روشی دقیق ارزیابی شود. در هر مورد، ماشین به گونه‌ای عمل نموده و تصمیم‌گیری می‌کند که می‌توان مستقیماً طراحی، برنامه‌نویسی و دانش تعبیه‌شده توسط اشخاص در ماشین را ردیابی کرد. پس در چنین مواردی اقدامات انسان، مستقیماً یا به دلیل توانایی در غلبه بر ماشین و به دست گرفتن راهبری آن، فرآیند را تعریف، هدایت و در نهایت کنترل می‌کند (Vladeck, 2014: 120). از این منظر، توسعه‌دهندگان را به‌ویژه در مقایسه با زیان‌دیده باید به عنوان شخصی در نظر گرفت که می‌تواند با کمترین هزینه از خطر مرتبط با فرآیند طراحی اجتناب کند و از آن‌جا که تلاش در راستای جلوگیری از بروز خسارت از جانب ایشان اصولاً موثر می‌باشد، می‌توان گفت که تلاش و حصول نتیجه در این موارد واقعیت‌هایی انکارناپذیر بوده و لذا با تشدید مسئولیت توسعه‌دهنده عدالت نیز محقق می‌شود (کریمیان، ۱۴۰۱: ۲۵). علاوه بر این، فراتر از اقدامات پیش‌گیرانه، توسعه‌دهندگان اغلب نهادهای تجاری هستند و بنابراین بازیگرانی هستند که بیشترین منافع را کسب می‌کنند. با توجه به پیامدهای نظریه اقتصادی، توسعه‌دهندگان باید تحت یک سازکار دارای انگیزه شوند که مشابه اقداماتی است که برای تولیدکنندگان سنتی اعمال می‌شود. از بند هفتاد و نه مقدمه دستورالعمل شماره ۱۶۸۹ مصوب ۲۰۲۴

اتحادیه اروپا که بیان داشته: «مناسب است یک شخص حقیقی یا حقوقی خاص که به عنوان ارائه‌دهنده تعریف می‌شود، مسئولیت عرضه در بازار یا به کار انداختن یک سیستم هوش مصنوعی پرخطر را بر عهده بگیرد، صرف نظر از این که آن شخص حقیقی یا حقوقی شخص طراح یا توسعه دهنده است»، می‌توان پی برد که مسئولیت توسعه دهنده تابع قاعده مسئولیت محض می‌باشد. هم‌چنین، بند دوازده ماده ۵۷ و بند نه ماده ۶۰ دستورالعمل مورد اشاره، ارائه‌دهنده را بر اساس قوانین اتحادیه و مسئولیت ملی قابل اجرا در قبال هر گونه آسیبی که در جریان آزمایش آن‌ها در شرایط دنیای واقعی ایجاد شود، مسئول دانسته است؛ بند نه ماده ۶۰ مقرر داشته: «ارائه‌دهنده یا ارائه‌دهنده احتمالی، طبق قوانین مربوط به مسئولیت اتحادیه و ملی، در قبال هرگونه خسارتی که در جریان آزمایش آن‌ها در شرایط دنیای واقعی ایجاد شود، مسئول خواهد بود».<sup>۱</sup>

از مطالب عنوان شده می‌توان پی برد که هوش مصنوعی صرف نظر از برخی موارد استثنا که باید به وسیله قانون‌گذار ممنوع اعلام شود،<sup>۲</sup> می‌تواند دارای خطر محدود یا پرخطر باشد. در نتیجه برای سیستم‌های هوش مصنوعی با خطر محدود، مسئولیت محض تنها در صورتی اعمال می‌شود که دادگاه ملی اثبات رابطه سببیت را برای مدعی بسیار دشوار بداند (بادینی و عباسی، ۱۳۹۵: ۶). تنها ایرادی که این نظریه دارد آن است که تفکیک و تشخیص مصادیق کم‌خطر و پرخطر دشوار است و حداقل باید ضابطه آن به وسیله قانون‌گذار ارائه شود. در بهینه‌سازی سطح فعالیت توسعه‌دهندگان، قابل پیش‌بینی بودن یک ریسک باید در نظر گرفته شود. مسئولیت محض توسعه‌دهندگان باید به ریسکی محدود شود که می‌توان بر اساس وضعیت متعارف پیش‌بینی کرد؛ زیرا هدف اصلی مسئولیت محض بازدارندگی است و در تشخیص قابلیت پیش‌بینی مستند به ماده ۵۲۱ قانون جدید مجازات اسلامی ضابطه

1. (79) It is appropriate that a specific natural or legal person, defined as the provider, takes responsibility for the placing on the market or the putting into service of a high-risk AI system, regardless of whether that natural or legal person is the person who designed or developed the system.

۲. در بخشی از بند بیست و پنج توصیه‌نامه یونسکو تحت عنوان اصل تناسب و پرهیز از آسیب بیان شده است که سیستم‌های هوش مصنوعی نباید از حدود آن‌چه که برای دستیابی به مقاصد و اهداف مشروع لازم است فراتر روند.

شخصی و نوعی هر دو مدنظر می‌باشد (صفایی و رحیمی، ۱۳۹۴: ۱۱۲) هر چند عملاً و در مقام اثبات صرف معیار نوعی برای احراز مسئولیت کفایت می‌نماید هر چند توسعه دهنده شخصا آن را پیش‌بینی نکرده باشد. اگر اثر بازدارنده جنبه فرعی داشته باشد، بار مسئولیتی که بیش از حد سنگین است، عمدتاً فعالیت‌های سودمند را محدود می‌کند. بدیهی است اشخاصی که به دنبال جبران خسارتی بیشتر از خسارت‌های تحت پوشش مسئولیت محض هستند، باید بر اساس قاعده مسئولیت مبتنی بر تقصیر به صورت مستقل علیه شخص مسئول اقدام نمایند. به نظر نگارندگان در ایران نیز پذیرش مسئولیت محض توسعه‌دهندگان مطابق با قواعد می‌باشد زیرا حتی با پذیرش شخصیت حقوقی هوش مصنوعی و بروز افعال پیش‌بینی ناپذیر این محصول، بر اساس قاعده استناد عرفی می‌توان مسئولیت را بر دوش عامل انسانی قرار داد و قاعده مزبور در قانون جدید مجازات اسلامی از جمله ماده ۵۲۶ مورد پذیرش قرار گرفته است.

## ۴-۲. مسئولیت مدیران سیستم‌های هوش مصنوعی

مدیران سیستم‌های هوش مصنوعی در موقعیتی هستند که می‌توانند در ایجاد خطر و یا کاهش آن نقش داشته باشند، و این مستلزم طرح قواعد کارآمد در باب مسئولیت است تا انگیزه‌هایی برای رفتار مناسب در اختیار آن‌ها قرار دهد. مسئله کلیدی این است که آیا مدیران باید مشمول مسئولیت محض باشند یا قواعد مبتنی بر تقصیر. ویژگی‌های سیستم‌های هوش مصنوعی و هم‌چنین زمینه‌ای که در آن مستقر شده‌اند، ارزیابی این‌که چه شخصی باید برای بهینه‌سازی سطح فعالیتش ترغیب شود را دشوار می‌کند. حداقل سه عامل در تعیین مسئولیت مدیران در عصر هوش مصنوعی باید در نظر گرفته شود.

### ۴-۲-۱. شدت خطر

هنگامی که مسئولیت محض و مسئولیت مبتنی بر تقصیر مقایسه می‌شود، یک سوال مهم از منظر بازدارندگی این است که آیا طرفین حتی پس از دستیابی به سطح بهینه احتیاط باید تشویق شوند تا سطح فعالیت خود را بیشتر کنترل کنند. پاسخ به این سوال، بسته به شدت زیان احتمالی مربوط به فعالیت یک مدیر، می‌تواند از موردی به مورد دیگر متفاوت باشد. اگر فعالیت یک طرف خاص بسیار خطرناک باشد، پاسخ

مثبت است، زیرا انگیزه اضافی برای رفتار صحیح می‌تواند خسارات اجتماعی را در حد قابل توجهی کاهش دهد. در مقابل، اگر یک سیستم فقط باعث زیان جدی شود، اعمال مسئولیت محض منطقی نیست، زیرا مانع از انجام فعالیت‌های مفید می‌شود (گندم‌کار و صالحی مازندرانی و حمیدی، ۱۴۰۰: ۲۴۳).

نتیجه این است که یک رویکرد مبتنی بر ریسک باید برای تعیین مسئولیت مدیران استفاده شود. در واقع، چنین رویکردی چیز جدیدی نیست. برای مثال، قانون‌گذاران نقش مهم کسانی را که دارایی‌های خطرناکی مانند حیوانات غیر اهلی، مواد خطرناک و وسایل نقلیه را نگهداری یا کنترل می‌کنند، تشخیص داده‌اند. مسئولیت محض این اشخاص را وادار می‌کند تا فعالیت‌های خود را بهینه کنند. منابع افزایش ریسک چارچوبی را برای تعیین مسئولیت مدنی یک شخص خاص در یک مدل مسئولیت محض یا به اصطلاح مسئولیت بدون تقصیر تعیین می‌کند (Kārklīņš, 2020: 165). در بیشتر موارد دیگر، مسئولیت مبتنی بر تقصیر به عنوان نظام پیش‌فرض برای مدیران عمل می‌کند. به نظر نگارندگان در ایران نیز نظریه تقصیر که در ماده ۱ قانون مسئولیت مدنی به عنوان یک قاعده عام پذیرفته شده است در این مورد قابل اعمال می‌باشد (صفایی و رحیمی، ۱۴۰۰: ۷۷) مگر آن‌که قانون‌گذار صراحتاً مبنای دیگری برای مسئولیت مدیران برگزیند.

با این حال، در زمینه هوش مصنوعی، سؤالی که باید بیشتر به آن پرداخته شود این است: چگونه می‌توانیم تعیین کنیم که آیا یک کاربرد خاص از هوش مصنوعی دارای درجه بالایی از خطر می‌باشد. قانون باید قلمرو اعمال بسیار خطرناک را در رابطه با هوش مصنوعی بیشتر روشن کند. یکی از مسائلی که در اینجا باید به آن اشاره کرد این است که اگرچه زیان واقعی برخی از فعالیت‌های فوق‌خطرناک را نمی‌توان به راحتی از طریق سیستم مسئولیت مدنی جبران کرد (مثلاً به صورت ضرر غیر مادی یا ضررهای اقتصادی خالص)، این بر ارزیابی از این‌که آیا چنین فعالیت‌هایی فوق‌خطرناک است یا خیر، تأثیری ندارد. بازدارندگی و جبران خسارت دو موضوع متفاوت هستند؛ لذا، می‌توان علاوه بر نظام مسئولیت مدنی، سازکارهای جایگزین را برای جبران خسارت ایجاد کرد.

## ۴-۲-۲. ناهمگونی عمل‌گرها

ناهمگونی مدیران مسائل دیگری را ایجاد می‌کند که باید در نظر گرفته شود. به‌طور سنتی، عملیات به این معنی است که یک طرف کنترل یک چیز را اعمال می‌کند و از مالکیت، استقرار یا استفاده از آن سود می‌برد. در عصر هوش مصنوعی، اشخاصی که این مورد برای آن‌ها اعمال می‌شود، مدیران نزدیک هستند. مدیران ممکن است از طریق تصمیم‌گیری خود در مورد مسائلی مانند زمان و نحوه عملکرد سیستم هوش مصنوعی باعث تضرر دیگری شوند، و بنابراین ارائه انگیزه به ایشان به منظور اعمال رفتار مناسب موجه است زیرا با گذشت زمان و بروز وقایع و رویدادهای جدید این احتمال به وجود می‌آید که ارائه‌دهنده راه حل رفع بن‌بست جدید، فردی غیر از مبتکران قبلی باشد (کریمیان، ۱۴۰۱: ۲۷). در شرایطی که سیستم‌های هوش مصنوعی را نمی‌توان به عنوان عامل خطرناک توصیف کرد، مسئولیت محض هیچ سود قابل توجهی ایجاد نمی‌کند و ممکن است فعالیت‌های ارزشمند را مهار کند (سلیمی و پارساپور، ۱۴۰۰: ۵۶). مسئولیت محض برای مدیران نزدیک را نمی‌توان به راحتی توجیه کرد مگر این‌که برنامه خاص هوش مصنوعی دارای سطح بالایی از خطر باشد.

در زمینه برنامه‌های کاربردی هوش مصنوعی، گروه‌های مختلفی وجود دارند که با ارائه خدمات پشتیبانی ضروری و مداوم، کنترل سیستم‌های هوش مصنوعی را اعمال و اطمینان حاصل می‌کنند که سیستم هوش مصنوعی قادر به تعامل مناسب با محیط است، به‌طوری که مدیران نزدیک نیز قادر به تعامل هستند، زیرا اشخاص ثالث می‌توانند بر عملکرد سیستم‌های هوش مصنوعی اعتماد کنند. به عنوان مثال، هنگامی که سامانه هوش مصنوعی در یک جاده در حال اجرا است، عملکرد آن به خدمات ارائه شده توسط اشخاص پشتیبان مختلف، مانند ارائه‌دهندگان داده، ارائه‌دهندگان خدمات مجازی، مدیران جاده‌ها، ارائه‌دهندگان سیستم‌های ماهواره‌ای ناوبری و غیره متکی است. آن‌ها بر عملکرد یک سیستم هوش مصنوعی نه تنها از طریق رفتار استراتژیک خود، بلکه مهم‌تر از آن، از طریق عملکرد خدمات خود تأثیر می‌گذارند. تشخیص این‌که آیا یک سرویس پشتیبانی خاص در همان سیستم برنامه واقعی هوش مصنوعی گنجانده شده است یا این‌که یک فرایند عملیاتی قابل تفکیک

است، در کل دشوار است.

نقش مدیران پشتیبانی از درک سنتی ما از آنچه عملیات نامیده می‌شود، فراتر رفته است، و مرز بین عملیات و ایجاد در حال تبدیل شدن است. ارتباط عملی خدمات پشتیبانی در زمینه بسیاری از برنامه‌های کاربردی هوش مصنوعی نشان می‌دهد که مدیران پشتیبان باید انگیزه‌هایی برای بهینه‌سازی فعالیت‌های خود به جای این که صرفاً به تدابیر پیش‌گیرانه اساسی اقدام نمایند، داشته باشند. با توجه به روشی که توسعه‌دهندگان و مدیران پشتیبان، فعالیت‌های یکدیگر را تکمیل می‌کنند، دشوار است که آن‌ها مشمول قوانین مسئولیت متفاوتی باشند. بر خلاف مدیران نزدیک، مدیران پشتیبان باید بدون توجه به ماهیت برنامه هوش مصنوعی مشمول مسئولیت محض باشند.

### ۳-۲-۴. سطح استقلال و درجه کنترل

یک سیستم هوش مصنوعی ممکن است به وسیله بیش از یک شخص اداره شود. در حالی که توسعه‌دهندگان، کنترل سیستم هوش مصنوعی را به صورت مداوم اعمال می‌کنند، مدیران نزدیک استفاده واقعی آن را تعیین می‌کنند. پیامدهای رفتار نادرست توسط مدیران نزدیک‌تر، به میزان کنترل آن‌ها بر روی یک برنامه کاربردی هوش مصنوعی بستگی دارد. معمولاً فرض بر این است که با داخلی‌سازی برخی از هزینه‌های حوادث توسط مدیران نزدیک‌تر (از طریق یک نظریه مبتنی بر تقصیر) می‌توانیم هزینه‌های اجتماعی را کاهش دهیم. این به مدیران نزدیک‌تر انگیزه‌هایی برای اقدامات احتیاطی کارآمد ارائه می‌دهد. با این حال، این سازکار تشویقی ممکن است آن گونه که انتظار می‌رود در عصر هوش مصنوعی موثر نباشد. با افزایش سطح استقلال، فضای باقی‌مانده برای تصمیم‌گیری مدیر بعدی محدود می‌شود. اگر یک برنامه هوش مصنوعی بسیار مستقل باشد، اقدامات مدیران نزدیک اصولاً زمینه‌ساز بروز حوادث شدید نمی‌گردد؛ بنابراین قرار دادن مدیران نزدیک، در این شرایط، در معرض یک بار مسئولیت عمده ممکن است واقعاً حوادث را کاهش ندهد. با این حال، مدیران نزدیک ممکن است در بروز حوادث ناشی از عدم نگهداری یا بروزرسانی برنامه‌های خود نقش داشته باشند.

در این راستا، به نظر می‌رسد باید اصل کلی مسئولیت مدیر را به جهت

مصالح عامه گوشزد نمود؛ زیرا با پذیرش این نظر زیان‌دیدگان در مطالبه خسارات وارده دچار سردرگمی نمی‌شوند و با مراجعه به مدیران ایشان نیز می‌توانند در قالب قائم‌مقامی به سایر اشخاص مرتبط از جمله توسعه‌دهندگان بر مبنای نظریه خطر مراجعه نمایند، زیرا به حکم قاعده من له الغنم فعلیه الغرم هر گاه کسی از عملی سود برد باید زیان‌های ناشی از آن را نیز تحمل نماید. در واقع از ظواهر امر مسئولیت مدیر ربات پیدا است و بار اثبات عدم مسئولیت بر عهده خواننده می‌باشد؛ با این شرط که خواهان اثبات کند خواننده به صورت انحصاری ربات را کنترل می‌کند (رجبی، ۱۳۹۸: ۴۵۸). پس در صورت نقص در تولید، خواننده یا مدیر باید ثابت کند که این نقص مربوط به تولیدکننده یا طراح می‌باشد. بدیهی است این نظر در صورتی منطقی است که ربات از سطح بالایی از استقلال برخوردار نباشد وگرنه اصل بر مسئولیت تولیدکننده یا توسعه‌دهنده می‌باشد. چنان‌که برخی در این خصوص به ماده ۱۲ کنوانسیون سازمان ملل متحد<sup>۱</sup> درباره استفاده از ارتباطات الکترونیکی در قراردادهای بین‌المللی استناد نموده‌اند که بیان می‌دارد، شخصی که رایانه را برنامه‌نویسی نموده است در نهایت باید مسئول پیامی باشد که به وسیله دستگاه تولید می‌شود و از آن افاده اصل نموده‌اند (محمدی و علی‌پور، ۱۴۰۰: ۱۲۰). بدیهی است که فرض اخیر مسئولیت را تنها بر عهده تولیدکننده می‌داند در حالی‌که از مفاد کنوانسیون مزبور قابل برداشت است که تولیدکننده تنها مسئول نبوده اما می‌تواند مسئول نهایی تلقی شود.

---

1. United Nation Convention

## برآمد

۱- توسعه‌دهندگان هوش مصنوعی تأثیر قابل توجهی بر بروز خطر و وقوع حوادث دارند و اغلب به دلیل توانایی مالی توان جبران خسارت را دارند. بر این اساس، آن‌ها باید مشمول مسئولیت محض باشند. چنین نظامی انگیزه لازم را برای مراقبت و فعالیت بهینه و هم‌چنین تسهیل مدیریت ریسک ایجاد می‌کند. در حقوق ایران نیز حتی در صورت پذیرش شخصیت حقوقی هوش مصنوعی، بر اساس نظریه استناد عرفی می‌توان مبادرت به توجیه مسئولیت محض توسعه‌دهنده و همه اشخاصی نمود که به مثابه ایشان در استفاده و به‌کارگیری هوش مصنوعی ایفای نقش نموده‌اند.

۲- همانند سایر مدیران نزدیک، مدیران نزدیک سیستم‌های هوش مصنوعی باید در اصل مشمول مسئولیت مبتنی بر تقصیر باشند مگر این‌که فعالیت‌های آن‌ها باعث خسارات اجتماعی قابل توجهی شود. رویکرد مبتنی بر ریسک که توسط قانون‌گذاران اتحادیه اروپا پیشنهاد شده است، گامی در جهت درست است، اما باید بیشتر توسعه داده شود تا موقعیت‌هایی را که مدیران نزدیک در آن کاملاً مسؤول هستند، تعریف کند. افزون بر این، موضوع مسئولیت مدیران غیر مستقیم را باید جدا از جایگاه مدیران نزدیک در نظر گرفت. نقش اساسی که توسط مدیران نخستین ایفا می‌شود نشان می‌دهد که آن‌ها به جای این‌که سیستم‌های هوش مصنوعی را «اداره» کنند، طرف «ایجادکننده» هستند. مسئولیت محض تضمین می‌کند که مدیران پشتیبان انگیزه‌های مناسبی دارند.

۳- از منظر اقتصادی لازم نیست یک طرف مناسب وجود داشته باشد و همه بازیگرانی که بر خطر وقوع حادثه تأثیر می‌گذارند باید تشویق شوند تا اقدامات پیش‌گیرانه را انجام دهند. کنترل اعمال شده توسط مدیران نزدیک، نیاز به ارائه انگیزه برای مدیران نخستین را از بین نمی‌برد و بر اساس مبنای عمومی مسئولیت مدنی یعنی نظریه تقصیر می‌توان اعمال مدیران را کنترل نمود. در صورت طرح مسئولیت شخص حقوقی نیز اشخاص مزبور را می‌توان بر اساس مبنای پیش‌گفته مسؤول تلقی نمود و لذا در صورتی که به ثالثی آسیب وارد شده باشد در ابتدا مالک و سپس، مدیر و یا سازنده آن می‌تواند مسؤول باشد. به بیان دیگر، مالک در هر حال در برابر زیان‌دیده مسؤول است و در برابر مالک سازنده و توسعه‌دهنده آن مسؤول می‌باشند.

## فهرست منابع

### الف. فارسی

- \* السان، مصطفی (۱۳۸۵)، «ابعاد حقوقی انتقال فناوری از طریق سرمایه‌گذاری خارجی»، پژوهش‌های حقوقی، دوره ۵، شماره ۹.
- \* بادی‌نی، حسن و عباسی، سمیه (۱۳۹۵)، «بررسی کارایی مسئولیت محض از دیدگاه تحلیل اقتصادی حقوق»، پژوهش‌های حقوق تطبیقی، دوره ۱۴، شماره ۱.
- \* بادی‌نی، حسن (۱۴۰۲)، فلسفه مسئولیت مدنی، چاپ پنجم، تهران: شرکت سهامی انتشار.
- \* بیگی حبیب‌آبادی، محمد و عیسایی تفرشی، محمد و شهبازی‌نیا، مرتضی (۱۳۹۹)، «تحلیل نظریه رکنیت مدیران شرکت‌های تجاری در حقوق آلمان، انگلیس و ایران»، پژوهش‌های حقوق تطبیقی، سال بیست و چهارم، شماره ۴.
- \* حکمت‌نیا، محمود و محمدی، مرتضی و واثقی، محسن (۱۳۹۸)، «مسئولیت مدنی ناشی از تولید ربات‌های مبتنی بر هوش مصنوعی خودمختار»، حقوق اسلامی، سال پانزدهم، شماره ۶۰.
- \* ذاکری‌نیا، حانیه (۱۴۰۲)، «ماهیت و مبنای مسئولیت مدنی ناشی از هوش مصنوعی در حقوق ایران و کشورهای اتحادیه اروپا»، پژوهش حقوق خصوصی، دوره ۲۰، شماره ۱.
- \* رجبی، عبدالله (۱۳۹۸)، «ضمان در هوش مصنوعی»، مطالعات حقوق تطبیقی، دوره ۱۰، شماره ۲.
- \* رحیمی، حبیب‌اله و طرف، فاطمه (۱۳۹۷)، «بررسی ویژگی‌های عدالت مطلوب در ماده اول قانون مسئولیت مدنی ایران در مقایسه با نظریه عدالت توزیعی جان راولز»، پژوهش حقوق خصوصی، سال هفتم، شماره ۲۵.
- \* سلیمی، فضا و پارساپور، محمدباقر (۱۴۰۰)، «مبنای مطلوب مسئولیت مدنی از دیدگاه تحلیل اقتصادی حقوق (با نگاهی به حقوق بازار)»، پژوهش حقوق خصوصی، دوره ۹، شماره ۳۵.
- \* صفایی، حسین و رحیمی، حبیب‌اله (۱۳۹۴)، مسئولیت مدنی (الزامات خارج از

قرارداد)، چاپ هشتم، تهران: سمت.

\* طاعتی سرشکه، محمد (۱۴۰۴)، «چالش‌های حقوق قراردادهای الکترونیکی در تجارت بین‌الملل»، مجله حقوقی دادگستری، دوره ۸۹، شماره ۱۳۲.

\* کریمیان، محمدوزین (۱۴۰۱)، «تحلیل رابطه ریاضی با حقوق: با تمرکز بر منابع حق و تکلیف و نظریه عدالت»، دیدگاه‌های حقوق قضائی، دوره ۲۷، شماره ۱۰۰.

\* گندم‌کار، رضاحسین و صالحی‌مازندرانی، محمد و حمیدی، محمدمهدی (۱۴۰۰)، «بررسی تطبیقی امکان وجود شخصیت حقوقی برای سامانه‌های هوشمند در فقه امامیه، حقوق ایران و حقوق غرب»، پژوهش تطبیقی حقوق اسلام و غرب، سال هشتم، شماره ۴.

\* محمدی، سمانه و علی‌پور، سحر (۱۴۰۰)، «بررسی مبانی و ارکان مسئولیت مدنی موتورهای جست‌وجوگر در سوءاستفاده از داده‌های کاربران»، حقوق فناوری‌های نوین، دوره ۲، شماره ۴.

\* محمدی، پژمان و طهماسب زاده، هادی و محمدکریمی، محمدجواد (۱۴۰۴)، «حمایت از دارندگان حقوق مرتبط»، مجله حقوقی دادگستری، دوره ۸۹، شماره ۱۳۲.

\* مشهدی‌زاده، علیرضا و قلی‌نیا، رضا (۱۴۰۱)، «مسئولیت مدنی کاربر در به‌کارگیری سیستم هوش مصنوعی در خودرو»، پژوهش‌های حقوقی، دوره ۲۱، شماره ۵۰.

\* مقدم، محمدجواد و حسینی، سید ابراهیم (۱۴۰۴)، «هوشمند سازی فرایندهای قضایی و عدالت کیفری؛ راهبردها و موانع»، مجله حقوقی دادگستری، دوره ۸۹، شماره ۱۳۱.

\* نظرپور، حمزه و ایزدی‌فرد، علی‌اکبر و محسنی‌دهکلانی، محمد (۱۴۰۰)، «بررسی فقهی حقوقی مسئولیت مدنی ربات»، فقه و اصول، سال پنجاه و سوم، شماره ۴.

\* ولی‌پور، علی و اسماعیلی، محسن (۱۴۰۰)، «امکان‌سنجی مسئولیت مدنی هوش مصنوعی عمومی ناشی از ایجاد ضرر در حقوق مدنی»، اندیشه حقوقی، سال دوم، شماره ۶.

## ب. انگلیسی

- \* Arcila, B. B. (2024). **AI liability in Europe: How does it complement risk regulation and deal with the problem of human oversight?** Computer Law & Security Review: The International Journal of Technology Law and Practice, 54.
- \* Bartlett, B. (2023). **The possibility of AI-induced medical manslaughter: Unexplainable decisions**, epistemic vices, and a new dimension of moral luck. Medical Law International, 23(3).
- \* Kingston, J. (2018). **Artificial intelligence and legal liability**. In Research and Development in Intelligent Systems XXXIII
- \* Vladeck, D. C. (2014). **Machines without principals: Liability rules and artificial intelligence**. Washington Law Review, 89, 117-150.
- \* Yeung, K. (2019). **Étude sur les incidences des technologies numériques avancées (dont l'intelligence artificielle) sur la notion de responsabilité**, sous l'angle des droits humains. Study of the Council of Europe, DGI. [Bilingual: English/French]

## پ. فرانسوی

- \* Mangematin, C. (2020). **Droit de la responsabilité civile et l'intelligence artificielle [Civil liability law and artificial intelligence]**. (Doctoral dissertation or monograph, in French). Latvian:
- \* Karklins, J. (2020). **Makslīgais intelekts un civiltiesiskā atbildība [Artificial intelligence and civil liability]**. Juridiska zinātne / Law, No. 13, pp.7-24. (in Latvian).
- \* Note on Yeung (2019). **This source is a Council of Europe study, officially published in both English and French**. It is listed under English sources but may also be consulted in its French version.