

مسئولیت مدنی و جبران خسارت در اثر تصمیم‌گیری الگوریتمی

و هوش مصنوعی در روابط خصوصی

(تاریخ دریافت ۱۴۰۴/۱۱/۱۵ تاریخ تصویب ۱۴۰۴/۱۲/۲۲)

مزدک نصوری^{۱*} آرین عرب‌نوری^۲

چکیده

تصمیم‌گیری الگوریتمی و سامانه‌های هوش مصنوعی به‌طور فزاینده‌ای در روابط خصوصی به کار گرفته می‌شوند؛ از اعتبارسنجی و اعطای تسهیلات بانکی و قیمت‌گذاری پویا تا گزینش نیروی انسانی، ارزیابی ریسک بیمه، مدیریت محتوای پلتفرم‌ها و حتی توصیه‌های پزشکی و حقوقی. این تحول، کارآمدی و سرعت را افزایش داده، اما در کنار آن، نوع تازه‌ای از زیان‌ها را نیز پدید آورده است: زیان‌های ناشی از خطاهای داده‌ای، سوگیری‌های آماری، سوگیری‌های منابع، امکان بروز توهم، عدم شفافیت مدل‌ها، خودپادگیری و تغییر رفتار سامانه در گذر زمان، و نیز دشواری انتساب و اثبات رابطه سببیت. پرسش اصلی این مقاله آن است که در روابط خصوصی، هنگامی که تصمیمی الگوریتمی سبب ورود خسارت می‌شود، چگونه باید مسئولیت مدنی را سامان داد و جبران خسارت را به‌نحوی منصفانه و کارآمد تضمین کرد. مقاله با تکیه بر مبانی حقوق مسئولیت مدنی، قواعد تقصیر، تسبیب، ضمان ناشی از عیب، مسئولیت ناشی از خطر، و نیز الزامات حرفه‌ای و قراردادی، تلاش می‌کند الگوهای انتساب مسئولیت را در زنجیره بازیگران متعدد تحلیل کند؛ از توسعه‌دهنده و عرضه‌کننده نرم‌افزار تا بهره‌بردار و کارفرما و ارائه‌دهنده داده. سپس راهکارهای جبرانی و پیشگیرانه، شامل معیارهای دقت و مراقبت متعارف، الزامات توضیح‌پذیری، مستندسازی، ممیزی، بیمه مسئولیت، و سازوکارهای قراردادی تنظیم ریسک ارائه می‌شود. نتیجه مقاله این است که راه‌حل پایدار، نه نفی فناوری و نه پذیرش بی‌قید آن، بلکه ترکیبی از معیارهای سخت‌گیرانه مراقبت حرفه‌ای، قواعد تسهیل اثبات، مسئولیت‌های چندلایه و توزیع ریسک در زنجیره عرضه است تا هم نوآوری تضعیف نشود و هم قربانی زیان، بدون جبران نماند.

واژگان کلیدی: مسئولیت مدنی، جبران خسارت، تصمیم‌گیری الگوریتمی، هوش مصنوعی، روابط خصوصی، تقصیر، سببیت، عیب محصول، توزیع ریسک

۱. دانشجوی دکتری حقوق عمومی، دانشگاه تهران، تهران، ایران (نویسنده مسئول) mazdak.nasouri@ut.ac.ir

۲. دانشجوی دکتری مهندسی کامپیوتر، دانشگاه تربیت مدرس، تهران، ایران arian_arabnouri@modares.ac.ir

مقدمه

گسترش خدمات دیجیتال و داده‌محور موجب شده است بسیاری از تصمیم‌هایی که پیش‌تر توسط انسان و بر پایه قضاوت فردی اتخاذ می‌شد، اکنون به سامانه‌های الگوریتمی و هوش مصنوعی واگذار شود. در روابط خصوصی، این تغییر بیش از هر جا در حوزه‌هایی نمود دارد که تصمیم باید سریع، انبوه و مبتنی بر داده باشد؛ مانند اعتباردهی، بیمه، استخدام، تبلیغات هدفمند، رتبه‌بندی کاربران، تشخیص تقلب، و تخصیص منابع. همین ویژگی‌ها سبب می‌شود خطای سامانه نیز به صورت انبوه و تکرارشونده رخ دهد و خسارت را در مقیاس گسترده ایجاد کند. در عصر حاضر، تصمیم‌گیری‌های مبتنی بر الگوریتم و سامانه‌های هوش مصنوعی از حاشیه به متن روابط حقوقی تنیده شده و به صورت فزاینده‌ای در هسته‌ی تعاملات خصوصی جای گرفته‌اند. دامنه شمول این فناوری‌ها از فرآیندهای سنتی اعتبارسنجی و اعطای تسهیلات اعتباری بانکی و قیمت‌گذاری‌های پویا فراتر رفته و تا مرزهای حیاتی‌گزینه‌ی نیروی انسانی، ارزیابی ریسک بیمه، کنترل و مدیریت محتوای پلتفرم‌های واسطه، و حتی ارائه مشاوره‌های تخصصی پزشکی و حقوقی گسترش یافته است. اگرچه این تحول پارادایمی، بدون شک منجر به ارتقای چشمگیر «کارآمدی» و «سرعت» در تبادلات حقوقی شده است، اما همزمان، منظومه‌ای نوین و پیچیده از زیان‌ها را بر جامعه تحمیل نموده است. این زیان‌ها که ماهیتی فراتر از خطاهای معمولی انسانی دارند، ریشه در ساختار ذاتی هوش مصنوعی دارند؛ زیان‌هایی که از خطاهای ساختاری در داده‌ها و بازتولید سوگیری‌های آماری و منابع آموزشی آغاز می‌شود، به پدیده‌های نوظهوری همچون «توهم»^۱ در خروجی‌های هوش مصنوعی ختم می‌گردد و در سایه «عدم شفافیت»^۲ مدل‌ها و قابلیت «خودیادگیری» آن‌ها که با گذر زمان

۱. امکان بروز توهم (hallucination) به این موضوع اشاره دارد که هوش مصنوعی اطلاعاتی تولید می‌کند که وجود خارجی ندارند. البته می‌توان خطر بروز آن را با تکنیک‌هایی مانند مهندسی پرامپت کاهش داد.

۲. علاوه بر عدم شفافیت که در مدل‌های منبع بسته هوش مصنوعی مانند gemini و chatgpt با آنها روبرو هستیم، مشکل جعبه سیاه (black box) یعنی اینکه تصمیمات بر چه اساسی گرفته شده‌اند، وجود دارد. البته مدل‌هایی وجود دارند که با ارایه منبع داده‌های مورد

موجب تغییر رفتار سامانه می‌شود، عمق می‌یابد؛ چالش‌هایی که در نهایت، پیچیده‌ترین موانع نظام حقوقی را در مسیر انتساب دقیق مسئولیت و اثبات رابطه سببیت میان تصمیم الگوریتمی و خسارت وارده پدید آورده است.

از منظر حقوقی، مسئله صرفاً جایگزینی ابزار نیست؛ بلکه تغییر در ماهیت تصمیم و نحوه شکل‌گیری آن است. تصمیم الگوریتمی به جای اتکای مستقیم به اراده و دانش تصمیم‌گیر انسانی، بر داده‌ها، مدل‌های آماری و قواعد محاسباتی متکی است. حتی در سامانه‌های یادگیرنده^۱، قواعد تصمیم ممکن است در طول زمان تغییر کند. در این وضعیت، پرسش‌هایی بنیادین مطرح می‌شود: اگر فردی به سبب ارزیابی اشتباه الگوریتم از دریافت خدمت محروم شود، یا قیمت‌گذاری تبعیض‌آمیز موجب زیان مالی او گردد، یا توصیه الگوریتمی به زیان جسمانی یا حیثیتی منتهی شود، مسئولیت بر عهده کیست. آیا باید از تقصیر انسانی سخن گفت یا از خطر ناشی از فناوری. چگونه می‌توان رابطه سببیت را در سامانه‌ای پیچیده اثبات کرد. و مهم‌تر، چگونه باید جبران خسارت را به نحوی تضمین کرد که قربانی زیان، در میان بازیگران متعدد و زنجیره پیچیده فناوری سرگردان نشود.

این مقاله در پی آن است که با حفظ چارچوب سنتی مسئولیت مدنی و بهره‌گیری از ظرفیت‌های آن، پاسخ‌های عملی ارائه کند. از یک سو، اصول عام مسئولیت مدنی همچنان نقطه اتکاست: زیان، فعل یا ترک فعل، رابطه سببیت، و عنصر انتساب. از سوی دیگر، ویژگی‌های هوش مصنوعی اقتضا می‌کند برخی مفاهیم بازخوانی شود؛ از جمله مفهوم تقصیر در حوزه فناوری پیشرفته، معیار مراقبت حرفه‌ای، نقش استانداردها و عرف فنی، و نیز راهکارهای تسهیل اثبات برای زیان‌دیده. همچنین، در کنار قواعد عام، قواعد خاصی مانند

استفاده سعی در کاهش این مشکل دارند، ولی این مسأله در الگوریتم‌های هوش مصنوعی سنتی مانند discriminative AI که از یادگیری عمیق استفاده می‌کنند به شدت وجود دارد.

۱. با استفاده از یادگیری تقویتی (reinforcement learning)

ضمان ناشی از عیب و مسئولیت تولیدکننده یا عرضه‌کننده می‌تواند در برخی مصادیق نقش کلیدی ایفا کند.

بیان مسئله

مسئله اصلی در مسئولیت مدنی ناشی از تصمیم‌گیری الگوریتمی و هوش مصنوعی در روابط خصوصی، تعارض میان دو نیاز است: نیاز به حمایت مؤثر از زیان‌دیده و نیاز به حفظ امکان نوآوری و توسعه فناوری. اگر معیار مسئولیت بیش از حد سخت‌گیرانه و نزدیک به مسئولیت مطلق طراحی شود، ممکن است هزینه‌های پیش‌بینی‌ناپذیر برای کسب‌وکارها و توسعه‌دهندگان ایجاد کند و انگیزه نوآوری را کاهش دهد. اگر معیار مسئولیت بیش از حد سهل‌گیرانه و مبتنی بر اثبات دشوار تقصیر باشد، قربانی زیان در عمل به جبران خسارت نمی‌رسد و عدالت ترمیمی و بازدارندگی از بین می‌رود.

در تصمیم‌گیری الگوریتمی، عوامل متعدد در ایجاد زیان دخالت دارد. گاه داده‌های ورودی ناقص یا آلوده است؛ گاه مدل به‌طور ساختاری سوگیر است^۱؛ گاه بهره‌بردار سامانه، آن را خارج از محدوده توصیه‌شده استفاده می‌کند؛ گاه تنظیمات و پارامترها به‌درستی انجام نمی‌شود؛ و گاه سامانه در طول زمان از محیط داده‌می‌آموزد و رفتاری متفاوت از نسخه اولیه بروز می‌دهد.^۲ در چنین فضایی، تعیین «فاعل» زیان‌آفرین دشوار می‌شود. حتی اگر بتوان فاعل را شناسایی کرد، اثبات اینکه تصمیم‌نهایی ناشی از تقصیر او بوده است، نیازمند دسترسی به مستندات فنی و توضیح مدل است؛ امری که معمولاً برای زیان‌دیده امکان‌پذیر نیست. از سوی دیگر، روابط خصوصی اغلب مبتنی بر قرارداد است؛ اما قرارداد همیشه راه‌حل نیست. بسیاری

۱. البته می‌توان سوگیری مدل را در فاز تست، اندازه‌گیری نمود. ولی مسأله سوگیری معمولاً در فرآیند آموزش و با توجه به داده صورت می‌پذیرد.

۲. البته موارد بیشتری هم وجود دارد. به‌عنوان مثال، حتی بهترین مدل‌ها هم احتمال تشخیص خطا که پایین هم نیست را دارند. همچنین، امکان نفوذ به این سیستم‌ها توسط بداندیشان وجود دارد. مضافاً اینکه امکان قرار گرفتن سیستم از حالت دسترس‌پذیر و نیز تداخل سیستم‌های هوش مصنوعی قابل‌تصور است.

از زیان‌ها نسبت به اشخاص ثالث رخ می‌دهد یا ناشی از شروط یک طرفه و عدم توازن قراردادی است. همچنین، در بسیاری از حوزه‌ها مانند پلتفرم‌ها، بانک‌ها و بیمه‌ها، طرف مقابل مصرف‌کننده یا کارجوست که توان چانه‌زنی ندارد. همچنین، حفظ حریم خصوصی کاربران و رعایت حقوق مالکیت و حقوق حفاظت از داده‌ها با نیازمندی‌های آموزش مدل‌ها دارای چالش موازنه است. لذا تکیه صرف بر آزادی قراردادی می‌تواند به بی‌پناهی زیان‌دیده منجر شود. بنابراین مسئله به صورت دقیق چنین صورت‌بندی می‌شود: با توجه به پیچیدگی فنی و چندبازیگری بودن سامانه‌های هوش مصنوعی، چه الگوی مسئولیت مدنی می‌تواند هم انتساب زیان را ممکن کند، هم بار اثبات را به نحوی منصفانه توزیع نماید، و هم جبران خسارت را به صورت مؤثر تضمین کند. افزون بر آن، باید روشن شود کدام ابزارها در حقوق خصوصی کارآمدتر است: تقصیر مبتنی بر معیار مراقبت حرفه‌ای، مسئولیت مبتنی بر خطر در برخی کاربری‌های پرریسک، ضمان عیب محصول، یا ترکیبی از اینها به همراه قواعد حمایتی و بیمه. در نهایت، مسئله تنها جبران خسارت گذشته نیست؛ بلکه پیشگیری از زیان‌های آینده نیز در قلب مسئولیت مدنی قرار دارد. اگر سامانه‌ها به گونه‌ای طراحی شوند که قابلیت ممیزی، توضیح‌پذیری و ردگیری تصمیم را فراهم کنند، هم کشف خطا سریع‌تر می‌شود و هم مسئولیت‌پذیری افزایش می‌یابد. پس باید رابطه میان الزامات فنی و معیارهای حقوقی نیز تبیین شود تا حقوق، به جای واکنش دیر هنگام، به تنظیم پیشینی ریسک کمک کند.

مرور ادبیات

ادبیات حقوقی مرتبط با این موضوع را می‌توان در چند محور دسته‌بندی کرد. محور نخست، آثار کلاسیک مسئولیت مدنی است که به مبانی تقصیر، تسبیب، نظریه خطر، و قواعد جبران خسارت می‌پردازد. این آثار اگرچه در زمانه پیشا-الگوریتمی نوشته شده‌اند، اما چارچوب مفهومی لازم را فراهم می‌کنند؛ از جمله بحث معیار رفتار متعارف، نقش عرف و تخصص، و امکان تعدیل معیار تقصیر بر اساس وضعیت فاعل و ماهیت فعالیت.

محور دوم، ادبیات نوین درباره فناوری‌های نو و مسئولیت مدنی است؛ مانند مسئولیت ناشی از محصولات معیوب، نرم‌افزار، خدمات دیجیتال و پلتفرم‌ها. در این حوزه، پژوهش‌ها نشان می‌دهد که تقسیم‌بندی سنتی میان «کالا» و «خدمت» و نیز تمایز میان «عیب ساخت» و «عیب طراحی» در نرم‌افزار پیچیده‌تر می‌شود. همچنین، مفهوم استانداردهای ایمنی و هشدارهای لازم، در نرم‌افزار و سامانه‌های یادگیرنده نیازمند بازتعریف است.

محور سوم، ادبیات خاص هوش مصنوعی و تصمیم‌گیری خودکار است که غالباً بر مسائل شفافیت، تبعیض، توضیح‌پذیری، و حساب‌دهی تمرکز دارد. در این ادبیات، از یک سو بر ضرورت سازوکارهای فنی مانند ثبت رخدادها، کنترل داده و ممیزی الگوریتم تأکید می‌شود و از سوی دیگر بر نیاز به قواعد حقوقی تسهیل‌کننده مسئولیت. برخی نویسندگان، مسئولیت مبتنی بر خطر را برای کاربری‌های پرریسک پیشنهاد می‌کنند تا زیان‌دیده از دشواری اثبات تقصیر رهایی یابد. در مقابل، برخی دیگر بر این باورند که مسئولیت سخت‌گیرانه می‌تواند نوآوری را محدود کند و بهتر است معیار تقصیر با در نظر گرفتن استانداردهای حرفه‌ای و تکالیف مراقبتی دقیق تقویت شود.

محور چهارم، ادبیات تطبیقی و مقررات‌گذاری است که به سیاست‌های حقوقی در اتحادیه اروپا، ایالات متحده آمریکا و برخی نظام‌های حقوقی دیگر می‌پردازد. در اروپا، مباحث پیرامون مقررات هوش مصنوعی، حفاظت داده، و جهت‌گیری به سوی ارزیابی ریسک و تکالیف پیشینی بسیار برجسته است. در ایالات متحده آمریکا، چارچوب مسئولیت مدنی غالباً در دل حقوق ایالات و دکترین‌های تقصیر و مسئولیت محصول پیگیری می‌شود و تمرکز بیشتری بر دعاوی خصوصی و رویه قضایی دیده می‌شود. مقایسه این دو رویکرد برای حقوق ایران آموزنده است؛ زیرا نشان می‌دهد چگونه می‌توان میان حمایت از زیان‌دیده و جلوگیری از بارگذاری افراطی بر نوآوری تعادل برقرار کرد.

در ادبیات فارسی نیز در سال های اخیر، نوشته هایی درباره مسئولیت مدنی در فضای مجازی، پلتفرم ها، داده های شخصی و چالش های هوش مصنوعی منتشر شده است. این آثار به ویژه بر مبانی فقهی و حقوقی ضمان، قاعده لاضرر، تسبیب، اتلاف، و نیز ظرفیت های قانون مسئولیت مدنی و قواعد عام مدنی تأکید دارند. با این حال، در بخشی از ادبیات فارسی، هنوز پیوند دقیق میان ویژگی های فنی سامانه های یادگیرنده و عناصر مسئولیت مدنی به صورت نظام مند تبیین نشده است. همچنین، موضوع توزیع مسئولیت میان توسعه دهنده، عرضه کننده، بهره بردار و دارنده داده، و نیز طراحی سازوکارهای اثبات و جبران خسارت، به پژوهش های دقیق تر نیاز دارد. این مقاله می کوشد با تلفیق این محورها، به یک چارچوب تحلیلی برسد که هم به مبانی حقوقی وفادار باشد و هم واقعیت فنی تصمیم گیری الگوریتمی را در نظر بگیرد؛ به ویژه با تأکید بر روابط خصوصی که در آن قراردادهای عدم توازن قدرت و نقش مصرف کننده و کارگر، اهمیت مضاعف دارد.

بخش اول: ماهیت تصمیم گیری الگوریتمی و جایگاه آن در روابط خصوصی

تصمیم گیری الگوریتمی، فرایندی است که در آن یک سامانه محاسباتی بر پایه مجموعه ای از داده ها و قواعد، خروجی ای تولید می کند که اثر تصمیم دارد. این خروجی ممکن است امتیاز، رتبه، توصیه یا تصمیم نهایی باشد. در بسیاری از کاربردها، تصمیم الگوریتمی به طور مستقیم یا غیرمستقیم بر حقوق و منافع اشخاص اثر می گذارد. نمونه روشن آن، نمره اعتبار است که بر امکان دریافت وام و نرخ بهره اثر دارد. نمونه دیگر، سامانه های گزینش اولیه در استخدام است که می تواند فرصت شغلی را از فردی سلب کند. در فضای پلتفرم ها نیز الگوریتم های رتبه بندی و حذف محتوا می توانند به حیثیت یا درآمد کاربران لطمه بزنند.

در روابط خصوصی، ویژگی مهم تصمیم گیری الگوریتمی، پیوند آن با قراردادهای و خدمات است. بانک، بیمه یا پلتفرم، به عنوان ارائه دهنده خدمت، غالباً قراردادی را پیشنهاد می کند که شروط آن از پیش تعیین شده است. در این قراردادهای، تصمیم گیری الگوریتمی گاه به عنوان

ابزار داخلی ارائه‌دهنده خدمت معرفی می‌شود و گاه اصلاً به صراحت ذکر نمی‌گردد. این وضعیت، از حیث شفافیت و رضایت آگاهانه اهمیت دارد. اگر طرف ضعیف قرارداد نداند تصمیمی که درباره او گرفته می‌شود بر چه مبنایی است، امکان اعتراض و اصلاح خطا کاهش می‌یابد و در نتیجه احتمال زیان افزایش پیدا می‌کند.

ویژگی دیگر، چندلایه بودن تصمیم است. گاه سامانه صرفاً پیشنهاد می‌دهد و انسان تصمیم نهایی را می‌گیرد. اما در عمل، انسان ممکن است به خروجی سامانه اعتماد کند و نقش او به تأیید صوری تقلیل یابد. در این حالت، از منظر مسئولیت مدنی باید بررسی کرد آیا تصمیم واقعاً انسانی است یا تصمیم سامانه به انسان «منتسب» شده است. اگر انسان صرفاً مهر تأیید باشد، استناد زیان به سامانه و بهره‌بردار آن معقول‌تر است.^۱

همچنین، یادگیری ماشین موجب می‌شود که رفتار سامانه در طول زمان تغییر کند.^۲ این تغییر ممکن است مطلوب باشد، اما از منظر حقوقی سبب می‌شود پیش‌بینی‌پذیری کاهش یابد. در مسئولیت مدنی، پیش‌بینی‌پذیری نقش کلیدی دارد، زیرا معیار تقصیر تا حدی با امکان پیش‌بینی ضرر گره خورده است. وقتی سامانه در محیط واقعی رفتار جدیدی نشان می‌دهد، پرسش این است که بهره‌بردار و توسعه‌دهنده تا چه حد باید این تغییرات را پیش‌بینی و کنترل کنند. پاسخ حقوقی می‌تواند بر مفهوم «نظارت مستمر» و «به‌روزرسانی مسئولانه» بنا شود؛ یعنی در سامانه‌های یادگیرنده، تکلیف مراقبت پایان نمی‌یابد و باید به صورت فرایندی و مداوم اعمال شود. در نهایت، تصمیم‌گیری الگوریتمی با مسئله داده پیوند عمیق دارد. بسیاری از خسارت‌ها ناشی از داده‌های غلط، ناقص یا تبعیض‌آمیز است. بنابراین، مسئولیت مدنی در این حوزه بدون تحلیل نقش داده و مالکیت و کنترل آن ناقص می‌ماند. باید دید چه کسی داده را

۱. در حال حاضر برای این مسأله مفهوم Human in the Loop پیشنهاد شده است. این اصطلاح ناظر بر آن است که انسان، تصمیم‌گیرنده نهایی می‌باشد و این مسئولیت انسان است که از نتایج هوش مصنوعی پیروی کند یا خروجی دیگری اتخاذ نماید.

۲. با یادگیری تقویتی می‌توان کاری کرد با استفاده از مکانیزم جریمه-پاداش، سیستم هوش مصنوعی به تدریج عملکرد خود را بهبود دهد. البته استفاده از آن اختیاری است. همچنین، با توجه به وابستگی خروجی هوش مصنوعی به ورودی آن با تغییر داده‌ها، عملکرد آن نیز تغییر می‌یابد.

گردآوری کرده، چه کسی آن را پاک سازی و برچسب گذاری کرده، و چه کسی آن را در تصمیم گیری به کار گرفته است. هر کدام از این حلقه ها می تواند منشأ تقصیر یا منشأ انتساب خطر باشد. نتیجه این بخش آن است که تصمیم گیری الگوریتمی در روابط خصوصی، نه یک پدیده تک عاملی، بلکه شبکه ای از عوامل انسانی و فنی است. بنابراین، مدل مسئولیت نیز باید شبکه ای و چندلایه باشد تا از یک سو امکان انتساب و جبران فراهم شود و از سوی دیگر بار مسئولیت به طور ناعادلانه بر یک بازیگر واحد تحمیل نگردد.

بخش دوم: بازخوانی عناصر مسئولیت مدنی در بستر هوش مصنوعی

برای تحقق مسئولیت مدنی، معمولاً وجود زیان، رابطه سببیت و عنصر انتساب ضروری دانسته می شود. هر چند نظام های حقوقی در بیان دقیق این عناصر تفاوت دارند، اما منطق کلی مشابه است. در مواجهه با هوش مصنوعی، هر یک از این عناصر با چالش های خاص روبه رو می شود. زیان در تصمیم گیری الگوریتمی، غالباً چندبعدی است. زیان مالی مانند محرومیت از وام، افزایش نرخ بیمه یا از دست دادن فرصت شغلی، آشکارتر است. اما زیان های غیرمالی نیز قابل توجه است؛ مانند لطمه به حیثیت در اثر برچسب گذاری نادرست به عنوان متقلب، یا آسیب روانی ناشی از تصمیم تبعیض آمیز. همچنین، در برخی کاربردها مانند سلامت، زیان بدنی نیز ممکن است رخ دهد. نکته مهم این است که زیان ممکن است نتیجه «تصمیم منفی» نباشد، بلکه نتیجه «تصمیم متفاوت» باشد. برای مثال، قیمت گذاری پویا ممکن است برای فردی قیمت بالاتری تعیین کند بدون آنکه او از قیمت های دیگر آگاه باشد. زیان در اینجا تفاوت تحمیلی و غیرمنصفانه است که تشخیص و اثبات آن دشوارتر است.

رابطه سببیت در سامانه های پیچیده، بیشترین چالش را دارد. گاه تصمیم الگوریتمی تنها یک عامل در کنار عوامل دیگر است. برای مثال، رد شدن درخواست وام ممکن است به امتیاز اعتباری، سیاست داخلی بانک، یا اطلاعات ناقص مشتری مربوط باشد. همچنین، خروجی الگوریتم ممکن است احتمالاتی باشد، نه حکم قطعی. از منظر حقوقی باید تعیین کرد آیا

خروجی الگوریتم سبب متعارف زیان بوده یا صرفاً یکی از زمینه‌ها. در اینجا، تحلیل سبب عرفی و معیار غالب بودن علت اهمیت می‌یابد. همچنین، در مواردی که تصمیم به صورت خودکار اجرا می‌شود و انسان مداخله ندارد، پیوند سببیت معمولاً روشن تر است.

عنصر انتساب نیز در هوش مصنوعی پیچیده می‌شود، زیرا فاعل مستقیم ممکن است ماشین باشد، اما ماشین شخص حقوقی نیست که بتوان اراده یا تقصیر به آن نسبت داد. در حقوق خصوصی، انتساب معمولاً به اشخاص حقیقی یا حقوقی صورت می‌گیرد که کنترل یا نفع از فعالیت دارند. بنابراین، باید دید کدام بازیگر بر سامانه کنترل مؤثر داشته و بهره اقتصادی می‌برد. در این تحلیل، مفهوم «کنترل» جایگاه محوری دارد. اگر توسعه‌دهنده صرفاً نرم‌افزار را فروخته و دیگر هیچ نقشی در عملیات ندارد، کنترل بیشتر در دست بهره‌بردار است. اما اگر سامانه به صورت خدمت مبتنی بر ابر ارائه می‌شود و به‌روزرسانی و تنظیمات در کنترل ارائه‌دهنده است، انتساب به ارائه‌دهنده خدمت تقویت می‌شود. همچنین، اگر داده‌های آموزش و تنظیم مدل توسط بهره‌بردار فراهم شده باشد، انتساب به او تقویت می‌گردد. در نهایت، عنصر تقصیر به عنوان یکی از مبانی مسئولیت، نیازمند بازتعریف معیار رفتار متعارف است. در حوزه هوش مصنوعی، رفتار متعارف با استانداردهای فنی، بهترین رویه‌ها، و عرف حرفه‌ای گره می‌خورد. از یک شرکت فناوری انتظار می‌رود آزمون‌های کافی انجام دهد، داده را پاک‌سازی کند، سوگیری را ارزیابی نماید، و سازوکار نظارت و پاسخ به خطا را طراحی کند. بنابراین، تقصیر می‌تواند در «طراحی»، «آموزش»، «پیاده‌سازی»، «نظارت» یا «اطلاع‌رسانی و هشدار» رخ دهد. این تنوع، به دادگاه و کارشناس امکان می‌دهد تقصیر را دقیق‌تر تحلیل کند، اما هم‌زمان کار زیان‌دیده برای اثبات را سخت‌تر می‌کند. همین دشواری، مبنای طرح قواعد تسهیل اثبات در بخش‌های بعدی خواهد بود.

بخش سوم: الگوهای مسئولیت در زنجیره بازیگران تصمیم گیری الگوریتمی

در یک سامانه هوش مصنوعی که در روابط خصوصی به کار می رود، معمولاً چند نقش اصلی وجود دارد: توسعه دهنده یا سازنده مدل، عرضه کننده نرم افزار یا خدمت، یکپارچه ساز یا پیمانکار پیاده سازی، بهره بردار یا کارفرما، ارائه دهنده داده یا کار گزار داده، و گاه شخص ثالث مانند شرکت ممیز یا ارائه دهنده زیر ساخت. زیان ممکن است نتیجه خطای یک حلقه یا ترکیب چند حلقه باشد. بنابراین، نظام مسئولیت باید امکان انتساب چند گانه و مسئولیت تضامنی یا نسبی را بررسی کند. مسئولیت بهره بردار یا ارائه دهنده خدمت در روابط خصوصی معمولاً پررنگ ترین است، زیرا او تصمیم را در برابر مشتری یا طرف قرارداد اعمال کرده و منفعت مستقیم می برد. حتی اگر سامانه توسط شرکت دیگری ساخته شده باشد، بهره بردار با انتخاب ابزار، تعیین سیاست ها، و اجرای تصمیم، نقش مؤثر دارد. در بسیاری از موارد، معیار تقصیر بهره بردار با تکلیف «انتخاب و نظارت» سنجیده می شود. اگر بهره بردار بدون ارزیابی کافی، سامانه ای را به کار گیرد که برای حوزه مورد نظر مناسب نیست یا نرخ خطای بالایی دارد، می توان تقصیر را احراز کرد. همچنین، اگر بهره بردار سازوکار اعتراض و اصلاح خطا را فراهم نکند، یا تصمیم را بدون بررسی انسانی در موارد حساس اعمال کند، معیار مراقبت نقض می شود. مسئولیت توسعه دهنده و عرضه کننده نرم افزار نیز در قالب مسئولیت ناشی از عیب یا تقصیر طراحی قابل طرح است. اگر الگوریتم به گونه ای طراحی شده باشد که در شرایط متعارف استفاده، خطای قابل پیش بینی ایجاد کند، یا اگر هشدارهای لازم درباره محدودیت ها ارائه نشده باشد، می توان مسئولیت را مطرح کرد. تفاوت مهم این است که توسعه دهنده گاه رابطه قراردادی مستقیم با زیان دیده ندارد. در اینجا، مبانی مسئولیت خارج از قرارداد و قواعد ضمان ناشی از عیب یا نظریه تقصیر نسبت به ثالث اهمیت می یابد.

یکپارچه ساز یا پیمانکار پیاده سازی نیز ممکن است مسئول باشد، زیرا بسیاری از خطاها در مرحله استقرار رخ می دهد: تنظیم نادرست پارامترها، اتصال غلط داده ها، یا تغییرات ناسازگار با

مدل. در اینجا مسئولیت او غالباً قراردادی در برابر بهره‌بردار است، اما ممکن است در برابر زیان‌دیده نیز مسئولیت خارج از قرارداد پیدا کند، اگر رفتار او سبب مستقیم زیان باشد.

ارائه‌دهنده داده یا کارگزار داده نیز در برخی موارد نقش کلیدی دارد. اگر داده‌های عرضه‌شده، ناقص یا نادرست باشد و بر اساس آن تصمیم زیان‌بار اتخاذ شود، می‌توان از تقصیر سخن گفت. البته تعیین معیار مراقبت در این حوزه نیازمند توجه به نقش و توانایی ارائه‌دهنده داده است. اگر او صرفاً انتقال‌دهنده باشد، مسئولیت محدودتر است. اما اگر او داده را پردازش و برچسب‌گذاری و آماده‌سازی کرده و آن را با ادعای کیفیت مشخص عرضه کرده باشد، مسئولیت افزایش می‌یابد. به‌طور کلی، در زنجیره بازیگران، می‌توان دو رویکرد برای توزیع مسئولیت تصور کرد. رویکرد نخست، تمرکز مسئولیت بر بهره‌بردار و ارائه‌دهنده خدمت است، زیرا او در نقطه تماس با زیان‌دیده قرار دارد و دسترسی به جبران را آسان‌تر می‌کند. سپس بهره‌بردار می‌تواند با رجوع به توسعه‌دهنده یا پیمانکار، هزینه را در زنجیره منتقل کند. این رویکرد به حمایت مؤثر از زیان‌دیده کمک می‌کند. رویکرد دوم، توزیع مسئولیت بر اساس سهم واقعی هر بازیگر در ایجاد خطر یا خطا است، که از منظر عدالت تخصیصی جذاب است اما ممکن است دعوا را پیچیده و طولانی کند. راه‌حل میانه می‌تواند این باشد که در برابر زیان‌دیده، مسئولیت اصلی بر عهده بهره‌بردار باشد و در روابط داخلی، نظام رجوع و تقسیم مسئولیت با معیارهای فنی و قراردادی تنظیم گردد. در روابط خصوصی، نباید از نقش قرارداد غفلت کرد. بسیاری از بازیگران با قراردادهای متوالی به هم متصل‌اند. قرارداد می‌تواند تکالیف ایمنی، استانداردهای عملکرد، تعهد به بروزرسانی، و مسئولیت‌های جبرانی را تعیین کند. اما قرارداد نمی‌تواند همیشه مسئولیت در برابر زیان‌دیده را حذف کند، به‌ویژه اگر زیان‌دیده مصرف‌کننده یا شخص ثالث باشد یا شروط قراردادی غیرمنصفانه باشد. بنابراین، تحلیل مسئولیت باید همزمان دو سطح را ببیند: سطح بیرونی که حمایت از زیان‌دیده است، و سطح درونی که توزیع ریسک میان بازیگران زنجیره عرضه است.

بخش چهارم: معیار تقصیر و مراقبت متعارف در سامانه های هوش مصنوعی

در مسئولیت مبتنی بر تقصیر، پرسش اصلی این است که رفتار مسئول چگونه باید باشد و آیا از آن معیار عدول کرده است یا خیر. در هوش مصنوعی، معیار مراقبت متعارف از چند منبع تغذیه می کند: عرف حرفه ای، استانداردهای فنی، رویه های صنعتی، و گاه مقررات تخصصی. بنابراین، قاضی برای تشخیص تقصیر ناگزیر است به دانش کارشناسی تکیه کند، اما این تکیه باید در چارچوب معیارهای حقوقی هدایت شود تا به نسبی گرایی و عدم قطعیت نینجامد.

تقصیر در مرحله طراحی می تواند شامل انتخاب الگوریتم نامناسب برای حوزه حساس، عدم پیش بینی سناریوهای خطا، یا عدم لحاظ کردن مخاطرات سوگیری باشد. برای مثال، به کارگیری مدلی که برای داده های یک جامعه دیگر آموزش دیده و بدون سازگارسازی در جامعه جدید استفاده می شود، می تواند رفتار نامتعارف باشد. همچنین، اگر سامانه در حوزه ای به کار می رود که تصمیم آن آثار مهم بر حقوق اشخاص دارد، انتظار می رود آزمون های جامع و ارزیابی ریسک انجام شود. این انتظار با مفهوم «سطح مراقبت بالاتر در فعالیت های پرخطر» قابل توجیه است.

تقصیر در مرحله داده، شامل گردآوری و پاک سازی ناکافی، عدم کنترل کیفیت، و استفاده از داده های غیرمرتبط یا ناقض حریم خصوصی است. بسیاری از خطاهای الگوریتمی در واقع خطاهای داده ای است. بنابراین، تکلیف مراقبت در داده به اندازه تکلیف مراقبت در کدنویسی اهمیت دارد. اگر بانک یا پلتفرم از منابع داده ای استفاده کند که به طور متعارف غیرقابل اعتماد است، یا داده های قدیمی و نادرست را بدون سازوکار به روزرسانی به کار گیرد، مسئولیت او قابل تصور است.

تقصیر در مرحله پیاده سازی و بهره برداری، شامل تنظیمات نادرست، استفاده خارج از محدوده توصیه شده، و اتکای کورکورانه به خروجی سامانه است. در اینجا مفهوم «نظارت انسانی معنادار» اهمیت می یابد. هر جا خروجی سامانه می تواند به زیان جدی منجر شود، بهره بردار باید

امکان بازبینی و مداخله انسانی را پیش‌بینی کند. البته این به معنی رد کامل خودکارسازی نیست؛ بلکه به معنی تعیین حوزه‌هایی است که نیازمند کنترل مضاعف است.

تقصیر در مرحله اطلاع‌رسانی و هشدار نیز قابل طرح است. عرضه‌کننده سامانه باید محدودیت‌ها، نرخ خطا، شرایط مناسب استفاده و مخاطرات شناخته‌شده را به‌طور قابل فهم بیان کند. بهره‌بردار نیز باید طرف قرارداد را از ماهیت تصمیم‌گیری خودکار و امکان اعتراض و اصلاح آگاه سازد. اگر زیان ناشی از بی‌اطلاعی معقول زیان‌دیده باشد، مسئولیت تقویت می‌شود.

نکته مهم آن است که معیار مراقبت در هوش مصنوعی ثابت نیست و با پیشرفت فناوری تغییر می‌کند. کاری که امروز به عنوان بهترین رویه شناخته می‌شود، ممکن است در چند سال آینده حد اقل استاندارد تلقی شود. بنابراین، نظام مسئولیت باید پویا باشد و از طریق معیار عرف حرفه‌ای و استانداردهای به‌روز، خود را تطبیق دهد. در عین حال، باید مراقب بود که این پویایی به بی‌ثباتی حقوقی منجر نشود. راهکار آن، توجه به معیارهای کلی مانند ارزیابی ریسک، مستندسازی تصمیمات فنی، و امکان ممیزی است که در هر دوره قابل تطبیق با فناوری موجود است.

بخش پنجم: چالش‌های اثبات و راهکارهای تسهیل آن در دعاوی خصوصی

حتی اگر مبانی مسئولیت روشن باشد، در عمل موفقیت دعاوی مسئولیت مدنی به مسئله اثبات وابسته است. در پرونده‌های ناشی از تصمیم‌گیری الگوریتمی، زیان‌دیده معمولاً با سه مانع روبه‌روست: عدم دسترسی به داده و مدل، پیچیدگی فنی، و ابهام در رابطه سببیت. این موانع موجب می‌شود بار اثبات سنتی که بر دوش خواهان است، در عمل به عدم جبران بینجامد.

نخستین مانع، عدم شفافیت سامانه است. بسیاری از مدل‌ها مانند شبکه‌های عصبی عمیق، توضیح ساده‌ای از چرایی تصمیم ارائه نمی‌کنند. علاوه بر آن، شرکت‌ها ممکن است مدل و

داده را به عنوان اسرار تجاری تلقی کنند و از افشای آن خودداری نمایند. نتیجه این می شود که زیان دیده نمی تواند ثابت کند تصمیم بر پایه داده غلط یا سوگیر بوده است. در چنین وضعی، حقوق می تواند به جای الزام به افشای کامل اسرار تجاری، سازوکارهای متوازن طراحی کند؛ مانند ارائه توضیح قابل فهم از عوامل مؤثر، یا دسترسی کنترل شده برای کارشناس بی طرف. هدف این است که امکان بررسی فراهم شود بدون آنکه حقوق مالکیت فکری بی جهت آسیب ببیند.

مانع دوم، پیچیدگی زنجیره ای است. زیان ممکن است حاصل ترکیب خطاهای چند عامل باشد. برای مثال، داده توسط یک شرکت گردآوری شده، مدل توسط شرکت دیگر ساخته شده، و بهره بردار با تنظیمات خاصی آن را به کار گرفته است. زیان دیده در این شبکه نمی داند باید علیه چه کسی طرح دعوا کند و چگونه نقش هر کدام را اثبات نماید. راهکار حمایتی در اینجا می تواند تمرکز بر بازیگر نزدیک تر به زیان دیده باشد، یعنی ارائه دهنده خدمت، و سپس اعطای حق رجوع به او در برابر سایر عوامل. این روش، هزینه دادرسی را برای زیان دیده کاهش می دهد و جبران را واقعی تر می کند.

مانع سوم، ابهام سببیت است. در تصمیم های مبتنی بر احتمال، ممکن است گفته شود حتی بدون الگوریتم نیز همان تصمیم اتخاذ می شد. یا ممکن است شرکت ادعا کند تصمیم نتیجه مجموعه عوامل بوده است. برای مواجهه با این وضعیت، می توان از معیارهای حقوقی مانند «سببیت متعارف»، «کفایت سبب»، و «غلبه احتمال» استفاده کرد. همچنین، می توان به سمت قواعدی رفت که در آن، اگر خواهان بتواند قرائن قوی بر خطا یا نقص سامانه ارائه دهد، بار اثبات به خواننده منتقل شود تا او نشان دهد که مراقبت لازم را انجام داده یا سببیت وجود نداشته است. این جابه جایی بار اثبات در حوزه هایی که اطلاعات در دست خواننده است، منطقی و منصفانه است.

علاوه بر این، باید نقش مستندسازی و ثبت رخدادها را جدی گرفت. اگر بهره‌بردار و عرضه‌کننده مکلف باشند تصمیمات الگوریتمی مهم را ثبت کنند، مسیر اثبات هموارتر می‌شود. چنین تکلیفی می‌تواند در قالب تکلیف قراردادی، تکلیف حرفه‌ای یا حتی شرط ضمنی در ارائه خدمات مطرح گردد. وقتی ثبت و ممیزی وجود دارد، اختلافات سریع‌تر حل می‌شود و مسئولیت‌پذیری افزایش می‌یابد. در نهایت، باید به نقش کارشناسی توجه کرد. پرونده‌های هوش مصنوعی بدون کارشناس قابل حل نیست، اما کارشناسی نیز باید هدفمند باشد. سؤالات کارشناسی باید بر معیارهای حقوقی متمرکز شود: آیا سامانه مطابق استانداردهای متعارف حوزه طراحی و آزمون شده است. آیا داده کنترل کیفیت داشته است. آیا هشدارها کافی بوده است. آیا بهره‌بردار از محدوده استفاده مجاز عدول کرده است. با این چارچوب، کارشناسی به جای تبدیل شدن به داوری فنی نامحدود، در خدمت تشخیص تقصیر و انتساب قرار می‌گیرد.

بخش ششم: مبانی مسئولیت مبتنی بر عیب و خطر در کاربردهای پرریسک

در کنار مسئولیت مبتنی بر تقصیر، دو مبنای دیگر نیز در هوش مصنوعی مطرح می‌شود: مسئولیت ناشی از عیب و مسئولیت مبتنی بر خطر. هر کدام مزایا و محدودیت‌های خود را دارد و انتخاب میان آنها به نوع کاربری و شدت ریسک بستگی دارد. مسئولیت ناشی از عیب، به‌ویژه در حوزه‌هایی قابل طرح است که سامانه به عنوان «محصول» یا جزء محصول عرضه می‌شود. اگر نرم‌افزار یا مدل، نقصی داشته باشد که امنیت و قابلیت اعتماد متعارف را کاهش دهد و در نتیجه زیان وارد کند، می‌توان مسئولیت را مطرح کرد. در این چارچوب، تمرکز بر کیفیت و ایمنی است نه بر رفتار مقصرانه. مزیت این رویکرد آن است که زیان‌دیده کمتر درگیر اثبات تقصیر می‌شود. اما دشواری آن در هوش مصنوعی این است که عیب همیشه به معنای «اشکال فنی واضح» نیست. گاه سامانه دقیقاً همان‌گونه که طراحی شده عمل کرده، اما طراحی برای آن حوزه نامناسب بوده است. بنابراین، مفهوم عیب باید شامل عیب طراحی و عیب هشدار نیز باشد، نه فقط عیب ساخت.

مسئولیت مبتنی بر خطر معمولاً در فعالیت‌هایی مطرح می‌شود که ذاتاً خطرناک است و حتی با رعایت مراقبت نیز احتمال زیان وجود دارد. در هوش مصنوعی، برخی کاربردها چنین ویژگی‌ای دارد؛ مانند کاربردهای پزشکی، حمل‌ونقل خودکار، یا سامانه‌های تصمیم‌گیری که اثر مستقیم بر حق حیات، سلامت یا آزادی اشخاص دارد. هرچند موضوع این مقاله روابط خصوصی است، اما همین روابط نیز می‌تواند در حوزه سلامت یا خدمات حساس رخ دهد. در این موارد، می‌توان استدلال کرد که بهره‌بردار یا ارائه‌دهنده خدمت، به دلیل ایجاد و بهره‌برداری از یک خطر غیرمتعارف، باید مسئولیت سخت‌گیرانه‌تری بپذیرد. مزیت آن حمایت قوی از زیان‌دیده است. اما ایراد آن این است که اگر بدون مرزبندی دقیق اعمال شود، می‌تواند دامنه‌ای بسیار گسترده پیدا کند و هزینه‌های غیرقابل مدیریت ایجاد نماید.

راهکار متوازن این است که مسئولیت مبتنی بر خطر به کاربری‌های مشخص و پرریسک محدود شود و در سایر موارد، مسئولیت مبتنی بر تقصیر تقویت شده به کار رود. تقویت شده یعنی معیار مراقبت بالا، تکلیف مستندسازی، تکلیف آزمون و ممیزی، و قواعد تسهیل اثبات. در نتیجه، برای اکثر روابط خصوصی مانند اعتباردهی یا تبلیغات، تقصیر تقویت شده می‌تواند کفایت کند، اما برای تصمیم‌هایی که احتمال زیان سنگین و غیرقابل جبران دارد، مسئولیت سخت‌گیرانه‌تر قابل دفاع است. همچنین، باید توجه کرد که حتی در نظام تقصیر، مفهوم «فرض تقصیر» یا «اماره تقصیر» می‌تواند نقش مشابهی با مسئولیت مبتنی بر خطر داشته باشد، بدون آنکه به مسئولیت مطلق تبدیل شود. اگر ثابت شود سامانه فاقد کنترل‌های متعارف بوده یا نرخ خطای غیرقابل قبول داشته است، می‌توان تقصیر را مفروض گرفت مگر اینکه خواننده خلاف آن را اثبات کند. این روش از نظر سیاست حقوقی، تعادل بهتری میان حمایت و نوآوری برقرار می‌کند.

بخش هفتم: جبران خسارت و شیوه‌های ارزیابی زیان در تصمیم‌گیری الگوریتمی

پس از احراز مسئولیت، مسئله جبران خسارت مطرح می‌شود. در زیان‌های ناشی از تصمیم‌گیری الگوریتمی، ارزیابی خسارت می‌تواند پیچیده باشد، زیرا گاهی زیان به شکل «فرصت از دست رفته» است. برای مثال، فردی به سبب ارزیابی نادرست از شغل محروم شده و نمی‌توان به سادگی تعیین کرد اگر استخدام می‌شد چه درآمدی داشت. یا فردی به سبب امتیاز اعتباری اشتباه، وام با نرخ بالاتر گرفته است و باید اختلاف هزینه محاسبه شود. همچنین، در زیان‌های حیثیتی یا تبعیض‌آمیز، تعیین میزان خسارت معنوی نیازمند دقت است.

در خسارت مالی مستقیم، اصل بر جبران کامل است؛ یعنی بازگرداندن زیان‌دیده به وضعیتی که بدون وقوع زیان داشت. این جبران می‌تواند شامل پرداخت تفاوت هزینه، بازپرداخت مبالغ اضافی، یا جبران درآمد از دست رفته باشد. در خسارت ناشی از فرصت از دست رفته، می‌توان از معیار احتمال استفاده کرد. یعنی دادگاه می‌تواند احتمال تحقق منفعت را ارزیابی و بر همان اساس خسارت را تقویم کند. این روش در بسیاری از نظام‌های حقوقی پذیرفته شده و می‌تواند در پرونده‌های الگوریتمی نیز کارآمد باشد.

در خسارت معنوی، باید به ماهیت تصمیم الگوریتمی توجه کرد. اگر تصمیم موجب برچسب‌گذاری ناروا، حذف حساب کاربری، یا انتشار پیامدهای منفی شده باشد، آسیب حیثیتی و روانی جدی است. در اینجا جبران صرفاً مالی ممکن است کافی نباشد. جبران می‌تواند شامل اعاده حیثیت، اصلاح داده، حذف برچسب، و ارائه توضیح رسمی نیز باشد. در روابط خصوصی، این نوع جبران‌های غیرمالی اهمیت دارد، زیرا می‌تواند آثار زیان را واقعاً کاهش دهد. بنابراین، دادگاه می‌تواند در کنار خسارت مالی، حکم به انجام اقدام اصلاحی نیز بدهد، به‌ویژه اگر زیان ناشی از داده غلط یا تصمیم قابل اصلاح باشد.

یکی از بحث های مهم، خسارت جمعی و گسترده است. تصمیم الگوریتمی می تواند گروه بزرگی را متضرر کند، مانند قیمت گذاری ناعادلانه برای دسته ای از کاربران. در این حالت، رسیدگی های فردی ممکن است ناکارآمد باشد. هرچند چارچوب های دعوای گروهی در برخی نظام ها توسعه یافته، در هر نظام حقوقی باید راه حل هایی برای کاهش هزینه پیگیری زیان دیدگان در نظر گرفت. حتی در نبود سازوکار رسمی دعوای گروهی، می توان از نقش نهادهای حمایتی، انجمن ها، یا مکانیسم های حل اختلاف جمعی بهره گرفت تا جبران خسارت عملی تر شود.

در نهایت، باید میان جبران و پیشگیری، پیوند برقرار کرد. اگر جبران خسارت تنها به پرداخت پول محدود شود، شرکت ها ممکن است آن را هزینه کسب و کار تلقی کنند. اما اگر جبران با الزام به اصلاح سامانه، ممیزی، و تغییر رویه همراه شود، اثر بازدارندگی افزایش می یابد.^۱ در حوزه تصمیم گیری الگوریتمی، بازدارندگی اهمیت ویژه دارد؛ زیرا خطاهای سیستماتیک می تواند به صورت مستمر تکرار شود. بنابراین، جبران خسارت باید به گونه ای طراحی شود که انگیزه اصلاح ساختاری را ایجاد نماید.

بخش هشتم: نقش قرارداد، شروط محدودکننده مسئولیت و توازن در روابط خصوصی

روابط خصوصی معمولاً بر قرارداد استوار است. ارائه دهندگان خدمات دیجیتال غالباً شروطی را در قراردادهای استاندارد درج می کنند که مسئولیت را محدود یا نفی می کند. در حوزه هوش مصنوعی، این شروط ممکن است شامل سلب مسئولیت از خطاهای الگوریتمی، محدود کردن جبران به مبلغی ناچیز، یا انتقال ریسک به کاربر باشد. از منظر عدالت قراردادی، باید بررسی کرد تا چه حد چنین شروطی معتبر و منصفانه است.

۱. در بسیاری از موارد، اصلاح سامانه هوش مصنوعی می تواند منجر به بیش برآزش داده ها و تضعیف عملکرد مدل شود. در نتیجه، این تصحیح نیاز فوری و قوی به دخالت و نظارت کارشناسان خبره هوش مصنوعی دارد.

اصل آزادی قرارداد اقتضا می کند طرفین بتوانند حدود مسئولیت را تعیین کنند، اما این اصل مطلق نیست. نخست، شرطی که به طور کلی مسئولیت ناشی از تقصیر سنگین یا رفتار غیرمتعارف را ساقط کند، با مبانی نظم عمومی و عدالت سازگار نیست. دوم، در قراردادهای الحاقی، عدم توازن قدرت چانه زنی می تواند شرط را غیرمنصفانه کند. سوم، اگر زیان دیده شخص ثالث باشد، شرط قراردادی اساساً نمی تواند حقوق او را زایل کند. بنابراین، در بسیاری از موارد، شروط سلب مسئولیت در برابر زیان دیده کارآمد نیست و تنها در روابط داخلی میان شرکت ها اثر دارد.

در روابط کسب و کار با کسب و کار، قرارداد می تواند ابزار مهم مدیریت ریسک باشد. طرفین می توانند سطح خدمت، معیارهای دقت، تعهد به پاسخ گویی، روش ممیزی، و سازوکار جبران را به طور دقیق تعیین کنند. همچنین، می توانند بیمه مسئولیت را شرط کنند و هزینه آن را در قیمت لحاظ نمایند. این تنظیم قراردادی، از یک سو شفافیت ایجاد می کند و از سوی دیگر اختلافات را کاهش می دهد.

در روابط کسب و کار با مصرف کننده یا کارجو، قرارداد باید با الزامات حمایت از طرف ضعیف متوازن شود. شرطی که تصمیم گیری الگوریتمی را پنهان کند یا حق اعتراض را سلب نماید، از منظر انصاف قابل دفاع نیست. حتی اگر شرط از نظر صوری پذیرفته شده باشد، رضایت واقعی بدون اطلاع و درک کافی حاصل نمی شود. بنابراین، در این حوزه باید بر تکلیف اطلاع رسانی، حق دسترسی به توضیح، و حق اصلاح تأکید کرد. این حقوق، هم به کاهش زیان کمک می کند و هم مسیر مسئولیت را روشن تر می سازد.

نکته دیگر، مسئولیت قراردادی در برابر اجرای نادرست تعهد است. اگر ارائه دهنده خدمت بر اساس قرارداد متعهد به ارائه خدمت عادلانه و دقیق باشد، تصمیم الگوریتمی نادرست می تواند نقض تعهد باشد. مزیت مسئولیت قراردادی این است که در برخی نظام ها بار اثبات ساده تر است و نیازی به اثبات تقصیر نیست، بلکه کافی است عدم اجرای صحیح تعهد ثابت شود. در

روابط خصوصی، ترکیب مسئولیت قراردادی و قهری می‌تواند پوشش جامع‌تری ایجاد کند. زیان‌دیده می‌تواند بر اساس نزدیک‌ترین مبنا که حمایت بیشتری می‌دهد اقدام کند، و دادگاه نیز می‌تواند با توجه به اوضاع و احوال، مبنای مناسب را اعمال نماید.

بخش نهم: بیمه، صندوق‌های جبرانی و سازوکارهای توزیع ریسک

یکی از راه‌های تضمین جبران خسارت در حوزه‌های نو و پرریسک، استفاده از سازوکارهای مالی مانند بیمه مسئولیت و صندوق‌های جبرانی است. در تصمیم‌گیری الگوریتمی، به سبب پیچیدگی اثبات و امکان خسارت گسترده، بیمه می‌تواند نقش مهمی داشته باشد. اگر ارائه‌دهنده خدمت یا بهره‌بردار بیمه مسئولیت داشته باشد، زیان‌دیده با احتمال بیشتری به جبران می‌رسد و ریسک مالی برای شرکت نیز قابل مدیریت می‌شود.

بیمه در این حوزه با چالش ارزیابی ریسک روبه‌روست. شرکت بیمه باید بداند سامانه چه میزان احتمال خطا دارد و چه نوع زیان‌هایی ممکن است رخ دهد. این امر موجب می‌شود بیمه‌گر نیز خواهان شفافیت، ممیزی و استانداردهای فنی باشد. در نتیجه، بیمه می‌تواند اثر تنظیم‌گر غیرمستقیم داشته باشد و شرکت‌ها را به رعایت استانداردهای مراقبت وادار کند. به بیان دیگر، بیمه تنها ابزار جبران نیست، بلکه ابزار پیشگیری نیز هست.

در برخی حوزه‌ها که خسارت می‌تواند بسیار گسترده باشد یا شناسایی مسئول دشوار است، صندوق‌های جبرانی نیز قابل تصور است. صندوق می‌تواند از محل حق عضویت شرکت‌های فعال در یک صنعت یا از محل درصدی از درآمد آنها تأمین شود و در مواردی که اثبات مسئولیت دشوار است، به زیان‌دیدگان کمک کند. البته طراحی صندوق نیازمند قواعد دقیق است تا به سوءاستفاده و بی‌انضباطی نینجامد. صندوق باید مکمل مسئولیت باشد، نه جایگزین آن. یعنی پس از پرداخت به زیان‌دیده، صندوق بتواند به مسئول واقعی رجوع کند یا معیارهای اصلاحی را اعمال نماید.

همچنین، قراردادهای داخلی در زنجیره عرضه می‌تواند توزیع ریسک را سامان دهد. برای مثال، توسعه‌دهنده می‌تواند تعهد کند که سامانه مطابق استانداردهای مشخص طراحی شده و در صورت عیب طراحی، خسارت بهره‌بردار را جبران نماید. بهره‌بردار نیز می‌تواند تعهد کند که سامانه را مطابق دستورالعمل استفاده کند و در صورت استفاده خارج از محدوده، مسئولیت را بپذیرد. این توافقات اگرچه در برابر زیان‌دیده اثر مستقیم ندارد، اما در نهایت هزینه‌ها را به حلقه‌ای منتقل می‌کند که کنترل بیشتری بر منشأ خطر دارد. چنین انتقالی، کارآمدی اقتصادی مسئولیت را افزایش می‌دهد.

بخش دهم: پیشنهادهای سیاستی و حقوقی برای نظام مسئولیت مدنی در روابط خصوصی

بر اساس تحلیل‌های پیشین، می‌توان مجموعه‌ای از پیشنهادها را برای سامان‌دهی مسئولیت مدنی و جبران خسارت در اثر تصمیم‌گیری الگوریتمی ارائه کرد. این پیشنهادها با فرض حفظ ساختار کلی حقوق خصوصی و استفاده از ظرفیت قواعد موجود تنظیم می‌شود.

نخست، تقویت معیار مراقبت حرفه‌ای برای بهره‌برداران و ارائه‌دهندگان خدمات تصمیم‌گیری خودکار ضروری است. هرچا تصمیم‌الگوریتمی اثر مهم بر منافع اشخاص دارد، بهره‌بردار باید ارزیابی ریسک، آزمون پیشینی، و سازوکار نظارت مستمر داشته باشد. این تکالیف می‌تواند به عنوان معیار تقصیر در رویه قضایی تثبیت شود. در نتیجه، شرکت نمی‌تواند با اتکا به پیچیدگی فناوری از مسئولیت بگریزد.

دوم، باید سازوکارهای تسهیل اثبات توسعه یابد. در پرونده‌هایی که اطلاعات فنی و داده در اختیار خواننده است، دادگاه می‌تواند با اماره‌ها و جابه‌جایی بار اثبات، تعادل ایجاد کند. اگر زیان‌دیده بتواند نشان دهد تصمیم کاملاً خودکار بوده یا قرائن قوی از خطا و تبعیض وجود دارد، خواننده باید کیفیت اقدامات مراقبتی خود را اثبات کند. این رویکرد با انصاف و منطق «کنترل اطلاعات» سازگار است.

سوم، باید شفافیت معقول و متناسب ایجاد شود. شفافیت به معنای افشای کامل کد و مدل نیست؛ بلکه به معنای ارائه توضیح قابل فهم از عوامل اصلی، امکان اعتراض، و ارائه اطلاعات لازم به کارشناس بی طرف است. چنین شفافیتی، هم حقوق زیان دیده را تقویت می کند و هم امکان دفاع منصفانه برای شرکت را فراهم می سازد.

چهارم، در کاربری های پرریسک، باید به سمت مسئولیت سخت گیرانه تر یا فرض تقصیر حرکت کرد. این کاربری ها باید به طور محدود و دقیق تعریف شود تا نوآوری، بی جهت محدود نشود. معیار تشخیص می تواند شدت زیان احتمالی، میزان وابستگی تصمیم به سامانه خودکار، و دشواری کنترل توسط زیان دیده باشد.

پنجم، نقش قرارداد باید از ابزار سلب مسئولیت به ابزار تنظیم مسئولانه ریسک تبدیل شود. قراردادهای استاندارد با مصرف کننده باید شامل اطلاع رسانی روشن درباره تصمیم گیری خودکار و سازوکار اعتراض باشد. شروط غیرمنصفانه سلب مسئولیت باید در پرتو اصول انصاف و حمایت از طرف ضعیف محدود گردد. در روابط کسب و کار با کسب و کار، قراردادهای فنی دقیق، ممیزی و بیمه می تواند به توزیع ریسک کمک کند.

ششم، توسعه بیمه مسئولیت و حتی ایجاد سازوکارهای جبرانی تکمیلی می تواند تضمین عملی جبران خسارت را تقویت کند. بیمه، هم بار مالی را قابل مدیریت می کند و هم شرکت ها را به رعایت استانداردها و ادار می سازد. در حوزه هایی که احتمال خسارت جمعی بالاست، باید به راه حل های جمعی نیز اندیشید تا هزینه پیگیری برای زیان دیدگان کاهش یابد.

در مجموع، این پیشنهادها در پی آن است که مسئولیت مدنی در مواجهه با هوش مصنوعی، نه فروپاشد و نه به مسئولیت مطلق تبدیل شود، بلکه با معیارهای جدید مراقبت و ابزارهای اثبات و جبران، کارآمدتر و منصفانه تر گردد.

نتیجه گیری

تصمیم گیری الگوریتمی و هوش مصنوعی در روابط خصوصی، صورت‌بندی سنتی مسئولیت مدنی را با پرسش‌های جدید مواجه کرده است. زیان‌ها می‌تواند مالی، معنوی یا بدنی باشد و در مقیاس گسترده رخ دهد. رابطه سببیت و انتساب در سامانه‌های پیچیده و چندبازیگری مبهم‌تر می‌شود و اثبات تقصیر برای زیان‌دیده دشوار است. با این حال، چارچوب حقوق مسئولیت مدنی همچنان ظرفیت پاسخ‌گویی دارد، مشروط بر آنکه معیارهای مراقبت حرفه‌ای تقویت شود، سازوکارهای تسهیل اثبات بکار گرفته شود، و توزیع ریسک در زنجیره عرضه به صورت چندلایه طراحی گردد. الگوی مطلوب، تمرکز بر مسئولیت بهره‌بردار و ارائه‌دهنده خدمت در برابر زیان‌دیده، همراه با امکان رجوع او به سایر بازیگران زنجیره است. در کنار آن، توسعه مسئولیت ناشی از عیب و استفاده محدود و هدفمند از مسئولیت مبتنی بر خطر یا فرض تقصیر در حوزه‌های پرریسک می‌تواند حمایت مؤثرتری ایجاد کند. همچنین، جبران خسارت باید از پرداخت پول فراتر رود و شامل اصلاح داده، اعاده حیثیت، و اقدامات پیشگیرانه باشد تا خطاهای سیستماتیک تکرار نشود. قراردادهای ابزار تنظیم مسئولانه ریسک باشند، نه پوششی برای سلب مسئولیت از طرف ضعیف. بیمه و سازوکارهای جبرانی تکمیلی نیز می‌تواند تضمین عملی جبران را تقویت کند. در نهایت، حقوق مسئولیت مدنی در برابر هوش مصنوعی، در حال ورود به مرحله‌ای است که در آن، معیارهای فنی و معیارهای حقوقی باید به زبان مشترک برسند. هرچه سامانه‌ها توضیح‌پذیرتر، قابل‌ممیزی‌تر و مستندسازی‌شده‌تر باشند، هم اختلافات کمتر می‌شود و هم عدالت جبرانی بهتر تحقق می‌یابد. هدف نهایی این است که فناوری در خدمت رفاه و کارآمدی باشد، نه اینکه با ایجاد «ابهام مسئولیت» راه فرار از جبران خسارت را هموار کند.

منابع و مآخذ

- ابراهیمی، محمد. مسئولیت مدنی و جبران خسارت. تهران: نشر میزان، چاپ پنجم، ۱۳۹۸.
- امامی، سیدحسن. حقوق مدنی. جلد اول و جلد دوم. تهران: انتشارات اسلامی، چاپ های متعدد.
- کاتوزیان، ناصر. الزام های خارج از قرارداد مسئولیت مدنی. تهران: شرکت سهامی انتشار، چاپ دهم، ۱۳۹۲.
- کاتوزیان، ناصر. قواعد عمومی قراردادها. جلد های مختلف. تهران: شرکت سهامی انتشار، چاپ های متعدد.
- صفایی، حسین. حقوق خانواده و قواعد عمومی مسئولیت. تهران: نشر میزان، چاپ های متعدد.
- لنگرودی، محمدجعفر جعفری. ترمینولوژی حقوق. تهران: گنج دانش، چاپ های متعدد.
- قانون مدنی جمهوری اسلامی ایران، مصوب ۱۳۰۷ با اصلاحات و الحاقات بعدی، تهران: روزنامه رسمی کشور.
- قانون مسئولیت مدنی، مصوب ۱۳۳۹، تهران: روزنامه رسمی کشور.
- معتمدنژاد، کاظم. حقوق ارتباطات. تهران: انتشارات سمت، چاپ های متعدد.
- قاسمی، سیدحسین. مسئولیت مدنی ناشی از فعل غیر و مبانی ضمان. قم: انتشارات بوستان کتاب، چاپ دوم، ۱۳۹۶.

- حسینی نژاد، محمد. حقوق فناوری اطلاعات و ارتباطات. تهران: انتشارات جنگل، چاپ سوم، ۱۳۹۹.
- باقری، علی. «مسئولیت مدنی ارائه‌دهندگان خدمات در فضای مجازی و چالش‌های اثبات». فصلنامه پژوهش‌های حقوق خصوصی. سال دوازدهم، شماره دوم، ۱۴۰۰.
- موسوی، فاطمه. «تحلیل حقوقی سوگیری الگوریتمی و آثار آن در مسئولیت مدنی». مجله مطالعات حقوق و فناوری. سال سوم، شماره اول، ۱۴۰۲.
- *European Commission. Proposal for a Regulation of the European Parliament and of the Council Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act). Brussels: Publications Office of the European Union, ۲۰۲۱.*
- *High-Level Expert Group on Artificial Intelligence. Ethics Guidelines for Trustworthy AI. Luxembourg: Publications Office of the European Union, ۲۰۱۹.*
- *OECD. OECD Principles on Artificial Intelligence. Paris: Organisation for Economic Co-operation and Development, ۲۰۱۹.*
- *Restatement (Second) of Torts. St. Paul: American Law Institute Publishers, ۱۹۶۵.*
- *Restatement (Third) of Torts: Products Liability. Philadelphia: American Law Institute, ۱۹۹۸.*
- *Goudkamp, James, and Donal Nolan (eds.). Winfield and Jolowicz on Tort. ۲۰th ed. London: Sweet & Maxwell, ۲۰۲۰.*
- *Howells, Geraint, and Thomas Wilhelmsson. EC Consumer Law. ۲nd ed. London: Routledge, ۲۰۲۱.*

- Pagallo, Ugo. *The Laws of Robots: Crimes, Contracts, and Torts*. Dordrecht: Springer, ۲۰۱۳.
- Surden, Harry. "Machine Learning and Law." *Washington Law Review*, Vol. ۸۹, No. ۱, ۲۰۱۴.
- Calo, Ryan. "Robotics and the Lessons of Cyberlaw." *California Law Review*, Vol. ۱۰۳, No. ۳, ۲۰۱۵.
- Zweig, Katharina A. *Invisible Decisions: Understanding How Algorithms Shape Our World*. Cambridge, MA: Harvard University Press, ۲۰۱۹.
- Casey, Bryan, and Mark A. Lemley. "A.I. and the Law." *Stanford Law Review*, Vol. ۷۱, ۲۰۱۹.
- Wachter, Sandra, Brent Mittelstadt, and Luciano Floridi. "Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation." *International Data Privacy Law*, Vol. ۷, No. ۲, ۲۰۱۷.
- Barfield, Woodrow, and Ugo Pagallo (eds.). *Research Handbook on the Law of Artificial Intelligence*. Cheltenham: Edward Elgar Publishing, ۲۰۱۸.
- Eidenmüller, Horst, and Gerhard Wagner (eds.). *Law and Artificial Intelligence*. Oxford: Oxford University Press, ۲۰۲۱.
- Schwartz, Paul M., and Daniel J. Solove. *Information Privacy Law*. ۷th ed. New York: Wolters Kluwer, ۲۰۲۲.
- Hacker, Philipp, et al. *Explainable AI: Legal and Ethical Challenges*. Berlin: Springer, ۲۰۲۰.