

فرا تر از خطای الگوریتم: بازاندیشی مسئولیت در سامانه‌های هوش مصنوعی نظامی (مطالعه موردی حادثه میناب)^۱

شیما عطار*

شیوا شکر الهی**

نوع مقاله: پژوهشی

چکیده

گسترش استفاده از سامانه‌های هوش مصنوعی در عملیات نظامی، پرسش‌های تازه‌ای درباره حدود مسئولیت حقوقی و اخلاقی در قبال خطاها و پیامدهای ناشی از این سامانه‌ها مطرح کرده است. در حادثه مدرسه شجره طیبه در میناب، که در جریان حمله نظامی ایالات متحده آمریکا و اسرائیل به ایران رخ داد، حدود ۱۶۸ دانش‌آموز جان خود را از دست دادند. در برخی گزارش‌ها و تحلیل‌های منتشر شده درباره این حادثه، احتمال نقش آفرینی سامانه‌های هوشمند در فرایند هدف‌گیری با بروز خطای مرتبط با آن‌ها مطرح شده است. مقاله حاضر با بهره‌گیری از روش توصیفی-تحلیلی و با تکیه بر مطالعه موردی حادثه مدرسه شجره طیبه، می‌کوشد بررسی کند که روایت انتساب این حادثه به «خطای هوش مصنوعی» تا چه اندازه با مبانی حقوق بین‌الملل بشردوستانه و چارچوب‌های نوین مسئولیت سازگار است و این روایت چه تأثیری بر انتساب مسئولیت به کنشگران انسانی و سازمانی دارد. چارچوب نظری پژوهش بر مفاهیمی چون شکاف مسئولیت، علیت شبکه‌ای، نظریه خلق خطر، کنترل انسانی معنادار و مسئولیت شبکه‌ای استوار است. تحلیل انجام‌شده نشان می‌دهد که حتی در فرض نقش آفرینی سامانه‌های هوشمند در فرایند هدف‌گیری، حادثه را نمی‌توان صرفاً به یک خطای الگوریتمی فروکاست، بلکه باید آن را نتیجه فرایندی پیچیده و چندلایه دانست که در آن طراحان سامانه، تأمین‌کنندگان داده، اپراتورها و سطوح مختلف فرماندهی هر یک می‌توانند در شکل‌گیری پیامد نهایی نقش داشته باشند. نوآوری این پژوهش در طرح مفهوم «مسئولیت شبکه‌ای» به‌عنوان چارچوبی تحلیلی برای فهم توزیع مسئولیت در سامانه‌های پیچیده مبتنی بر هوش مصنوعی و همچنین تأکید بر ضرورت حفظ و تقویت «کنترل انسانی معنادار» در کاربردهای نظامی این فناوری است. بر این اساس، استدلال می‌شود که تقلیل چنین حوادثی به «خطای هوش مصنوعی» بدون بررسی دقیق شواهد، ساختار تصمیم‌گیری و نقش کنشگران انسانی، می‌تواند تمرکز تحلیل را از مسئولیت‌های انسانی و سازمانی منحرف ساخته و فرایند پاسخگویی حقوقی را با چالش مواجه کند.

واژگان کلیدی

حادثه میناب، شکاف مسئولیت، علیت شبکه‌ای، کنترل انسانی معنادار، مخاصمات مسلحانه، مسئولیت بین‌المللی، هوش مصنوعی.

۱. این مقاله برگرفته از گزارش اظهارنظر کارشناسی در خصوص «ابعاد حقوقی حادثه میناب» است که در پژوهشگاه قوه قضاییه انجام شده است.

* استادیار حقوق، گروه آینده پژوهی قضایی، پژوهشکده قوه قضاییه، تهران، ایران (نویسنده مسئول).

Shimaattar@gmail.com

** دانشجوی دکتری حقوق عمومی، دانشکده حقوق و علوم سیاسی، دانشگاه علامه طباطبائی، تهران، ایران.

shivashokrollahi8@gmail.com

تاریخ پذیرش: ۱۴۰۵/۰۴/۲۳

تاریخ دریافت: ۱۴۰۵/۰۲/۲۰

مقدمه

در سال‌های اخیر، استفاده از سامانه‌های هوش مصنوعی در عرصه نظامی به سرعت گسترش یافته است. از پهبادهای خودکار و سامانه‌های شناسایی مبتنی بر داده‌های ماهواره‌ای گرفته تا ابزارهای هدف‌گیری هوشمند، همگی نشان می‌دهند که سامانه‌های هوش مصنوعی امروزه نقش و حضوری فزاینده در میدان‌های نبرد پیدا کرده‌اند. این فناوری‌ها اگرچه نوید افزایش دقت، کاهش تلفات انسانی و سرعت عمل را می‌دهند، اما همزمان پرسش‌های دشواری را درباره مسئولیت در قبال بروز خطاها یا آسیب‌ها پدید آورده‌اند. به بیان دیگر، هنگامی که یک سامانه هوشمند تصمیمی می‌گیرد که منجر به کشته شدن افراد غیرنظامی می‌شود، این پرسش مطرح است که چه کسی یا چه نهادی باید پاسخگو باشد: طراح سامانه، برنامه‌نویس الگوریتم، فرماندهی که حمله را مجاز دانسته، یا نهادی که از سامانه استفاده کرده است؟ (موسی‌زاده و آذریندار، ۱۴۰۳: ۵۶). این مسئله که در ادبیات حقوقی و اخلاقی با عنوان «شکاف مسئولیت»^۱ شناخته می‌شود، همچنان یکی از چالش‌های مهم این حوزه به‌شمار می‌رود.

گسترش سامانه‌های خودکار و هوشمند، مدل سنتی مسئولیت را که مبتنی بر ارتباط مستقیم میان عمل یک انسان و پیامد آن بود، با چالش‌هایی مواجه کرده است. در نظام حقوقی کنونی، مسئولیت کیفری و مدنی متوجه فرد یا افرادی است که مرتکب فعل یا ترک فعل شده‌اند. اما در مورد سامانه‌های هوشمند، گاه نمی‌توان به سادگی تشخیص داد که کدام حلقه از زنجیره طراحی، آموزش داده، نظارت و اجرا نقش تعیین‌کننده‌تری در وقوع حادثه داشته است (زمانیان و وطنی، ۱۴۰۴: ۱۱). از همین رو، شماری از صاحب‌نظران بر ضرورت «کنترل انسانی معنادار»^۲ تأکید کرده‌اند؛ به این معنا که اقدامات نظامی مرگبار نباید بدون نظارت مؤثر و آگاهانه انسان انجام شود (محمودی، انصاری مهباری و راعی دهقی، ۱۴۰۴: ۷). با این حال، این پرسش مطرح است که آیا صرف حضور یک اپراتور در کنار سامانه، که گاه تنها چند ثانیه فرصت دارد تصمیم‌نهایی را تأیید یا رد کند، می‌تواند به معنای «کنترل معنادار» باشد؟ یا در عمل، این اپراتور خود به بخشی از یک فرایند تصمیم‌گیری پیچیده تبدیل می‌شود و فشار زمانی و حجم داده‌ها توانایی قضاوت مستقل را محدود می‌کند؟ (Balis, O'Neill, 2021: 19-22).

در چنین بستری، حادثه مدرسه «شجره طیبه» در میناب به‌عنوان یک مطالعه موردی انتخاب شده است. در پی این حادثه، اگرچه گزارش رسمی، نقش سامانه‌های هوش مصنوعی را در فرایند هدف‌یابی تأیید نکرد، برخی گزارش‌های رسانه‌ای و تحلیل‌های کارشناسی احتمال استفاده از

1. Gap of responsibility

2. Meaningful human control

سامانه‌های هوشمند، خطای هدف‌یابی یا اتکا به داده‌های قدیمی را مطرح کردند. از این رو، هدف این پژوهش اثبات نقش هوش مصنوعی در حادثه میناب نیست، بلکه بررسی پیامدهای حقوقی و تحلیلی چنین روایت‌هایی در فرایند انتساب مسئولیت در حقوق بین‌الملل بشردوستانه است. از این منظر، ارزش مطالعه موردی میناب نه در اثبات نقش هوش مصنوعی، بلکه در فراهم ساختن بستری برای ارزیابی کفایت الگوهای سنتی انتساب مسئولیت در شرایط عدم قطعیت است.

بررسی روایت‌های منتشرشده درباره حادثه میناب این پرسش را مطرح می‌کند که آیا نسبت دادن حادثه به «خطای الگوریتم» ممکن است برخی حلقه‌های علی در فرایند تصمیم‌گیری را نادیده بگیرد و در نتیجه تحلیل مسئولیت را با ابهام مواجه سازد؟

این پژوهش بر آن است که آنچه در برخی روایت‌ها به‌عنوان «خطای الگوریتم» توصیف می‌شود، ممکن است حاصل مجموعه‌ای از تصمیم‌ها و کاستی‌های انسانی در مراحل مختلف باشد؛ از طراحی سامانه و انتخاب داده‌های آموزشی گرفته تا تعریف معیارهای «هدف نظامی»، تنظیم آستانه حساسیت سامانه، نظارت میدانی و تأیید نهایی حمله توسط سطوح فرماندهی. در چنین شرایطی، فروکاستن این فرایند پیچیده به تعبیر «خطای سامانه هوش مصنوعی» ممکن است مانع از بررسی کامل مسئولیت‌ها شود و در نتیجه روند پاسخگویی حقوقی را با دشواری روبه‌رو کند (انصاری مهباری و محمودی، ۱۴۰۱: ۶).

پژوهش حاضر ابتدا چارچوب نظری مسئولیت در سامانه‌های هوش مصنوعی نظامی را بر پایه مفاهیمی چون شکاف مسئولیت، علیت شبکه‌ای، نظریه خلق خطر و کنترل انسانی معنادار تبیین می‌کند. سپس با استفاده از حادثه میناب به‌عنوان مطالعه موردی، قابلیت این چارچوب در تحلیل انتساب مسئولیت ارزیابی شده و در نهایت، یافته‌ها در پرتو اصول حقوق بین‌الملل بشردوستانه تحلیل و جمع‌بندی می‌شود.

۱. مروری بر جایگاه فناوری هوش مصنوعی در فرایندهای نظامی

توسعه سریع سامانه‌های هوش مصنوعی در حوزه نظامی، به‌ویژه در زمینه شناسایی اهداف، تحلیل داده‌های اطلاعاتی و پشتیبانی از تصمیم‌گیری عملیاتی، موجب شده است که مسئولیت حقوقی ناشی از تصمیمات این سامانه‌ها به یکی از مهم‌ترین مباحث حقوق بین‌الملل بشردوستانه تبدیل شود. اگرچه این فناوری می‌تواند دقت و سرعت پردازش اطلاعات را افزایش دهد، استفاده از آن در تصمیم‌گیری درباره کاربرد نیروی قهری، پرسش‌های بنیادینی درباره مسئولیت دولت‌ها، فرماندهان، اپراتورها و توسعه‌دهندگان سامانه‌های هوشمند ایجاد کرده است. از همین رو، نهادهایی مانند کمیته بین‌المللی صلیب سرخ و سازمان ملل متحد بر ضرورت حفظ «کنترل انسانی

معنادار» بر سامانه‌های تسلیحاتی مبتنی بر هوش مصنوعی تأکید کرده‌اند (ICRC, 2024; United Nations, 2024).

برای تحلیل مسئولیت در این حوزه، ابتدا باید نقش هوش مصنوعی در چرخه عملیات نظامی، الگوهای مختلف حضور انسان در تصمیم‌گیری و مهم‌ترین مخاطرات ناشی از اتکای فزاینده به این سامانه‌ها تبیین شود.

۱-۱. نقش هوش مصنوعی در چرخه شناسایی، تحلیل داده و هدف‌گیری

در عملیات‌های نظامی، چرخه شناسایی و هدف‌گیری شامل گردآوری اطلاعات، تحلیل داده‌ها، ارزیابی اهداف و تصمیم‌گیری درباره حمله است. پیش‌تر این فرایند عمدتاً متکی بر نیروی انسانی بود، اما با ورود فناوری‌هایی مانند «یادگیری ماشین»^۱، «بینایی رایانه‌ای»^۲، و «تحلیل کلان‌داده»^۳، بخش قابل توجهی از این چرخه خودکار شده است (محمودی، انصاری مهیاری و راعی دهقی، ۱۴۰۴: ۹-۱۱).

سامانه‌های هوشمند می‌توانند حجم عظیمی از تصاویر ماهواره‌ای، داده‌های پهپادی و اطلاعات مخابراتی را در زمانی کوتاه پردازش کرده و اهداف احتمالی را برای بررسی به اپراتور انسانی پیشنهاد دهند. همچنین برخی سامانه‌ها قادرند اهداف را اولویت‌بندی کرده و گزینه‌های عملیاتی مناسب را پیشنهاد کنند و در برخی موارد نیز در قالب سامانه‌های کاملاً خودکار، بدون مداخله مستقیم انسان اقدام نمایند (موسی‌زاده و آذرپندار، ۱۴۰۳: ۱۱). با وجود این، کمیته بین‌المللی صلیب سرخ تأکید کرده است که استفاده از چنین سامانه‌هایی نباید جایگزین قضاوت انسانی در تصمیم‌گیری درباره استفاده از زور شود؛ زیرا مسئولیت رعایت اصول بنیادین حقوق بین‌الملل بشردوستانه همچنان بر عهده انسان باقی می‌ماند (ICRC, 2024).

در عین حال، استفاده از سامانه‌های یادگیرنده با چالش «ابهام الگوریتمی»^۴ همراه است؛ به این معنا که حتی طراحان سامانه نیز همواره قادر به تبیین نحوه رسیدن الگوریتم به یک نتیجه مشخص نیستند. در چنین شرایطی، انتساب مسئولیت حقوقی پیچیده‌تر شده و مستلزم بررسی مجموعه عوامل انسانی، فنی و سازمانی مؤثر در فرایند تصمیم‌گیری خواهد بود (Pfaff, 2019: 134-136).

۲-۱. الگوهای حضور انسان در تصمیم‌گیری

اسناد سازمان ملل متحد و کمیته بین‌المللی صلیب سرخ، حفظ «کنترل انسانی معنادار» را یکی از

1. Machine learning
2. Computer vision
3. Big data analysis
4. Algorithmic ambiguity

پیش شرط‌های اساسی مسئولیت‌پذیری حقوقی در سامانه‌های تسلیحاتی مبتنی بر هوش مصنوعی می‌دانند. هرچه میزان مداخله مؤثر انسان کاهش یابد، انتساب مسئولیت حقوقی نیز پیچیده‌تر خواهد شد (United Nations, 2024; ICRC, 2024).

برای مدیریت مخاطرات ناشی از خودکارسازی فرایندهای نظامی، پژوهشگران و نهادهای بین‌المللی سه الگوی اصلی برای حضور انسان در زنجیره تصمیم مبتنی بر هوش مصنوعی معرفی کرده‌اند. این الگوها که در اسناد مختلف با عناوینی چون «درون حلقه»^۱، «روی حلقه»^۲ و «خارج از حلقه»^۳ شناخته می‌شوند، در این پژوهش با تعابیر «کنترل مستقیم»، «نظارت غیرمستقیم» و «اختیار کامل به ماشین» صورت‌بندی می‌شوند (زمانیان و وطنی، ۱۴۰۴: ۱۰-۱۴).

در الگوی کنترل مستقیم، سامانه هوشمند نقش مشورتی یا پیشنهادی دارد و تمامی تصمیمات مهم، به‌ویژه تصمیم به استفاده از زور، باید به تأیید انسان برسد. اپراتور می‌تواند خروجی سامانه را بپذیرد، رد کند یا خواستار اطلاعات تکمیلی شود. این الگو در مقایسه با سایر الگوها، بیشترین سطح پاسخگویی انسانی را فراهم می‌کند، زیرا هیچ اقدامی بدون فرمان آگاهانه یک انسان انجام نمی‌شود. با این حال، در عمل، فشار زمانی میدان جنگ و حجم انبوه داده‌ها می‌تواند اپراتور را به پذیرش سریع و کم‌تأمل پیشنهاد ماشین سوق دهد. در چنین شرایطی، انسان عملاً به «تأیید‌کننده خودکار» تبدیل شده و «کنترل مستقیم» رنگ می‌بازد (Balis & O'Neill, 2022: 22-24).

در الگوی نظارت غیرمستقیم، سامانه مجاز است به‌طور خودکار علیه اهداف از پیش تعریف‌شده اقدام کند، اما یک اپراتور انسانی عملکرد آن را در پس‌زمینه رصد کرده و امکان مداخله یا غیرفعال‌سازی دارد. این الگو معمولاً برای اهداف متحرک یا موقعیت‌هایی به کار می‌رود که سرعت واکنش انسانی از توانایی لازم برای مقابله با تهدید کمتر است. چالش اصلی این مدل آن است که اپراتور به‌دلیل نظارت همزمان بر چندین سامانه، ممکن است دچار خستگی یا غفلت شود و متوجه خطای قریب‌الوقوع سامانه نگردد. تجربه نشان داده است که انسان‌ها در نظارت منفعل و طولانی‌مدت بر عملکرد ماشین‌ها، ضریب خطای بالایی دارند (انصاری مهباری و محمودی، ۱۴۰۱: ۱۱).

در الگوی اختیار کامل به ماشین، هیچ انسانی در میانه فرایند تصمیم و اقدام حضور مؤثر ندارد و سامانه به‌طور کامل خودگردان عمل می‌کند. این وضعیت شدیدترین شکل شکاف مسئولیت را ایجاد می‌کند؛ زیرا در لحظه تصمیم، کنشگر انسانی مستقیمی وجود ندارد و انتساب مسئولیت به مراحل طراحی یا آموزش سامانه نیز با دشواری روبه‌رو است. از همین رو، بسیاری از صاحب‌نظران حقوق

1. Inside the ring
2. on the ring
3. outside the ring

بین‌الملل بشردوستانه، استفاده از چنین سامانه‌هایی را به‌ویژه علیه اهداف انسانی، محل مناقشه دانسته و بر محدودسازی یا ممنوعیت آن تأکید کرده‌اند (خلیفه و اسماعیل تبار، ۱۴۰۲: ۳-۵). مطالعه موردی حادثه میناب نیز صرفاً در فرض استفاده از سامانه‌های هوشمند قابل تحلیل در چارچوب این الگوها است. از آنجا که تاکنون اطلاعات رسمی و قابل راستی‌آزمایی درباره نوع سامانه‌های به‌کاررفته و میزان مداخله انسان منتشر نشده، این پژوهش درصدد اثبات نقش هوش مصنوعی در حادثه نیست؛ بلکه از این چارچوب برای تحلیل حقوقی موقعیت‌های محتمل انتساب مسئولیت بهره می‌گیرد.

۱-۳. پیامدهای عملیاتی و ریسک‌های ناشی از اتکای فزاینده به سامانه‌های هوشمند

اتکای روزافزون به سامانه‌های هوشمند، افزون بر پیچیده‌تر شدن انتساب مسئولیت، مخاطرات عملیاتی متعددی نیز ایجاد می‌کند. یکی از مهم‌ترین این مخاطرات، پدیده «رانس مهارت»^۱ است؛ بدین معنا که وابستگی بیش از حد به خروجی سامانه‌های هوشمند، توان تحلیل و قضاوت مستقل اپراتورها را کاهش می‌دهد و در شرایط بروز خطا، تصمیم‌گیری صحیح را دشوار می‌سازد (Webb, 2021: 2115-2117).

مخاطره دیگر، «سوگیری اتوماسیون» است که در آن کاربران، حتی در صورت مشاهده نشانه‌های خطا، تمایل دارند خروجی سامانه را صحیح فرض کنند. این پدیده یکی از مهم‌ترین موانع تحقق کنترل انسانی معنادار محسوب می‌شود (Scharre, 2018; ICRC, 2024). همچنین، «سوگیری الگوریتمی»^۲ ناشی از داده‌های آموزشی نامتوازن یا ناکامل می‌تواند احتمال تشخیص نادرست اهداف را افزایش دهد و در نتیجه، ارزیابی مسئولیت حقوقی را با پیچیدگی بیشتری مواجه سازد (بنافی، ۱۴۰۲: ۱۴-۱۳).

مخاطره مهم دیگر، «تسریع غیرقابل کنترل» است. در یک رویارویی نظامی که هر دو طرف از سامانه‌های هوشمند با سرعت بالا استفاده می‌کنند، ممکن است زنجیره‌ای از واکنش‌های خودکار شکل گیرد که انسان قادر به دخالت مؤثر در آن نباشد. الگوریتم ممکن است تهدیدی را تشخیص داده و ظرف کسری از ثانیه پاسخ دهد؛ حتی اگر اپراتور انسانی متوجه خطا شود، زمان کافی برای توقف در اختیار نخواهد داشت (Chernavskikh, Palayer, 2025: 12-15). این وضعیت که در برخی تحلیل‌های راهبردی از آن با عنوان «جنگ سرد الگوریتمی»^۳ یاد شده، به‌ویژه زمانی خطرناک است که دو سوی درگیر از تفسیرهای متفاوتی از نشانه‌ها برخوردار باشند (Can &

1. Skill drift
2. Algorithmic bias
3. Algorithmic Cold War

(Barış, 2019: 33-36).

بر این اساس، تحلیل مسئولیت در سامانه‌های هوشمند مستلزم بررسی کل زنجیره تصمیم است؛ مبنایی که در بخش‌های بعدی برای تحلیل حقوقی مطالعه موردی میناب به کار گرفته می‌شود.

۲. چارچوب نظری مسئولیت در سامانه‌های هوش مصنوعی نظامی

برای تحلیل مسئولیت در عملیات‌های نظامی مبتنی بر هوش مصنوعی، اتکای صرف به الگوهای سنتی مسئولیت که بر رابطه مستقیم میان فاعل، فعل و نتیجه استوارند، کفایت نمی‌کند؛ زیرا در این سامانه‌ها تصمیم‌ها حاصل تعامل مجموعه‌ای از کنشگران انسانی، سامانه‌های فنی و ساختارهای سازمانی است. از این رو، این پژوهش چارچوب نظری خود را بر پایه پنج مفهوم «شکاف مسئولیت»، «علیت شبکه‌ای»، «نظریه خلق خطر»، «کنترل انسانی معنادار» و «مسئولیت شبکه‌ای» استوار می‌کند.

در این چارچوب، میان سه سطح مسئولیت تفکیک می‌شود: مسئولیت کیفری اشخاص، مسئولیت بین‌المللی دولت و مسئولیت ناشی از طراحی، توسعه و بهره‌برداری از سامانه‌های هوشمند. مفاهیم یادشده جایگزین قواعد سنتی مسئولیت نیستند، بلکه ابزارهایی برای تحلیل نحوه انتساب مسئولیت در سامانه‌های پیچیده به‌شمار می‌آیند.

۲-۱. مسئله «شکاف مسئولیت» در فناوری‌های خودکار

نخستین و بنیادی‌ترین مفهوم در ادبیات حقوقی و اخلاقی پیرامون فناوری‌های خودکار، «شکاف مسئولیت» است. این اصطلاح برای نخستین بار در اوایل دهه ۲۰۰۰ میلادی توسط فیلسوفانی چون آندریاس ماتیاس و رابرت اسپارو مطرح شد (Matthias, 2004: 178-180) و به وضعیتی اشاره دارد که در آن، به دلیل پیچیدگی و خودمختاری سامانه، دیگر نمی‌توان هیچ عامل انسانی مشخصی را به طور مستقیم مسئول پیامدهای زیان‌بار سامانه دانست، در حالی که خود سامانه نیز به دلیل فقدان اراده و آگاهی، واجد شرایط مسئولیت کیفری یا اخلاقی نیست (زمانیان و وطنی، ۱۴۰۴: ۱۶). بدین ترتیب، یک «شکاف» یا «خلأ» میان دو سر طیف مسئولیت پدید می‌آید؛ از یک سو عامل انسانی که نمی‌تواند به تنهایی پاسخگو باشد و از سوی دیگر عامل ماشینی که ذاتاً پاسخگو نیست (Sparrow, 2007: 63-66).

در سال‌های اخیر، ادبیات حقوق بین‌الملل بشردوستانه رویکرد محتاطانه‌تری نسبت به مفهوم «شکاف مسئولیت» اتخاذ کرده است. کمیته بین‌المللی صلیب سرخ تأکید می‌کند که افزایش سطح خودکارسازی نباید به ایجاد خلأ در پاسخگویی حقوقی منجر شود و استفاده از سامانه‌های هوشمند، مسئولیت تصمیم‌گیران انسانی را از بین نمی‌برد، بلکه ضرورت تبیین دقیق‌تر زنجیره

تصمیم‌گیری و کنترل انسانی را دوچندان می‌کند (ICRC, 2024). در حقوق کیفری کلاسیک، تحقق مسئولیت مستلزم احراز دو رکن اساسی است: «عنصر مادی» و «عنصر روانی». در مورد یک سامانه هوشمند نظامی که خود تصمیم به حمله می‌گیرد، به سختی می‌توان گفت کدام انسان مرتکب «رفتار» مجرمانه شده است. طراح الگوریتم ممکن است سال‌ها پیش کد را نوشته و پیش‌بینی رفتار سامانه در تمامی شرایط عملیاتی برای او ممکن نبوده باشد. تأمین‌کننده یا برچسب‌گذار داده‌های آموزشی شاید ناآگاهانه مجموعه‌ای محدود یا سوگیرانه فراهم کرده باشد. اپراتوری که تأیید نهایی را صادر کرده، ممکن است صرفاً بر اساس خروجی سامانه و تحت فشار زمانی عمل کرده باشد. حتی فرمانده‌ای که دستور استقرار سامانه را داده، الزاماً فرمان حمله به هدف خاصی را صادر نکرده است (Santoni, Filippo & Van den Hoven, 2018: 8-10).

در چنین شرایطی، علیت در طول زمان و میان کنشگران متعدد توزیع می‌شود و یافتن یک «فاعل انسانی واحد» که بتواند بار کامل مسئولیت را بر عهده گیرد، دشوار و گاه ناممکن می‌نماید. این وضعیت همان چیزی است که از آن به‌عنوان شکاف مسئولیت یاد می‌شود. با این حال، پرسش بنیادین آن است که آیا این شکاف، واقعیتی اجتناب‌ناپذیر و ذاتی فناوری‌های خودکار است، یا حاصل اصرار بر الگوی سنتی و فردمحور مسئولیت؟ شاید راه برون‌رفت از این بن‌بست نه در انکار مسئولیت، بلکه در بازتعریف آن نهفته باشد؛ مسئولیت نه به‌مثابه صفتی که به یک فرد منفرد تعلق می‌گیرد، بلکه به‌عنوان ویژگی یک شبکه از کنشگران انسانی است که در طراحی، آموزش، استقرار و نظارت بر سامانه نقش داشته‌اند. این رویکرد با تحلیل حادثه میناب نیز سازگارتر است. در این پژوهش نیز حادثه میناب صرفاً به‌عنوان مطالعه موردی برای ارزیابی چارچوب‌های انتساب مسئولیت بررسی می‌شود و نه به‌عنوان دلیلی بر اثبات نقش سامانه‌های هوش مصنوعی؛ زیرا تاکنون اطلاعات رسمی و قابل راستی‌آزمایی درباره نقش این سامانه‌ها منتشر نشده است (Amnesty International, 2025; Reuters, 2025).

۲-۲. مفهوم علیت شبکه‌ای در سامانه‌های پیچیده

برای عبور از دشواری‌های ناشی از شکاف مسئولیت، برخی نظریه‌پردازان به مفهوم «علیت شبکه‌ای» متوسل شده‌اند. این دیدگاه که ریشه در فلسفه علم و مطالعات علم و فناوری دارد، بر این نکته تأکید می‌کند که در سامانه‌های بسیار پیچیده مانند نیروگاه‌های هسته‌ای، شبکه‌های حمل‌ونقل هوشمند یا سامانه‌های تسلیحاتی خودگردان، دیگر نمی‌توان علیت را به‌صورت خطی و یک‌جهته از علت به معلول تصور کرد (Latour, 2005: 70-72).

در ادبیات جدید حقوق فناوری و هوش مصنوعی نیز مفهوم علیت شبکه‌ای به عنوان پاسخی به محدودیت الگوی سنتی علیت مورد توجه قرار گرفته است. بر اساس این رویکرد، تصمیمات اتخاذ شده توسط سامانه‌های مبتنی بر هوش مصنوعی، حاصل تعامل پیچیده میان الگوریتم، داده‌های آموزشی، زیرساخت‌های فنی، تنظیمات عملیاتی، مداخلات انسانی و ساختارهای سازمانی است؛ از این رو، انتساب نتیجه زیان‌بار به یک عامل منفرد، تصویر کاملی از منشأ واقعی حادثه ارائه نمی‌کند. به همین دلیل، بسیاری از پژوهشگران پیشنهاد می‌کنند که تحلیل مسئولیت در سامانه‌های هوشمند بر پایه «زنجیره تصمیم‌گیری» و «شبکه کنشگران» انجام شود، نه صرفاً عملکرد نهایی الگوریتم (Santoni, Filippo & Van den Hoven, 2018: 4).

در چارچوب علیت شبکه‌ای، نتیجه زیان‌بار نه محصول یک علت منفرد، بلکه حاصل برهم‌کنش یک شبکه درهم‌تنیده از عوامل انسانی و غیرانسانی است؛ شبکه‌ای که در آن هر کنشگر (انسان، سازمان، نرم‌افزار، سخت‌افزار، رویه اداری، مقررات و حتی شرایط محیطی) سهمی در شکل‌گیری نتیجه دارد، اما هیچ‌یک به تنهایی «علت تامه» حادثه محسوب نمی‌شود. از این منظر، تحلیل مسئولیت باید بر کل زنجیره تصمیم‌گیری و شبکه کنشگران متمرکز باشد، نه صرفاً بر عملکرد نهایی الگوریتم.

۳-۲. نظریه خلق خطر و نقش طراحان و بهره‌برداران

سومین مفهوم کلیدی که می‌تواند به پر کردن شکاف مسئولیت در سامانه‌های هوش مصنوعی نظامی کمک کند، «نظریه خلق خطر» است. این نظریه که ریشه در حقوق مدنی و به‌ویژه بحث «مسئولیت ناشی از فعالیت‌های خطرناک» دارد، بر این اصل استوار است که هر شخص حقیقی یا حقوقی که با رفتار خود، خطری بیش از حد معمول برای دیگران ایجاد کند، باید مسئول جبران خسارات ناشی از تحقق آن خطر باشد، حتی اگر نتوان اثبات کرد که رفتار او مستقیماً و به‌طور مشخص علت حادثه بوده است (کاتوزیان، ۱۳۹۹: ۲۳۴-۲۳۶).

در این چارچوب، هنگامی که یک دولت یا یک بنگاه نظامی و صنعتی تصمیم می‌گیرد از یک سامانه هوشمند و خودگردان در عملیات نظامی استفاده کند، در واقع دست به اقدامی می‌زند که به‌طور روشن واجد وصف خلق خطر است. این خطر شامل مجموعه‌ای از ریسک‌های خاص می‌شود: پیش‌بینی‌ناپذیری رفتار سامانه‌های یادگیرنده، امکان بروز سوگیری‌های الگوریتمی، وابستگی تصمیم‌انسانی به خروجی‌های ماشینی، و احتمال خطاهای فاجعه‌بار در شرایط پیچیده و پرتنش میدانی. این نوع خطرهای پیش از استقرار سامانه وجود نداشته یا دست‌کم به این شدت و کیفیت مطرح نبوده‌اند (Asaro, 2022: 212-215).

حال اگر چنین خطری محقق شود و به آسیب غیرنظامیان بینجامد، دولت یا بنگاه مذکور نمی‌تواند صرفاً با ارجاع به «خطای الگوریتم» از مسئولیت شانه خالی کند. تصمیم به طراحی، استقرار و به کارگیری سامانه، منشأ وضعیتی پرخطر بوده است که پیامدهای آن، هرچند در جزئیات غیرقابل پیش‌بینی، اما در کلیت قابل انتظار بوده‌اند. همین منطقی در فرض استفاده از سامانه‌های هوشمند در حادثه میناب نیز قابل اعمال است؛ بدین معنا که تحلیل مسئولیت باید فراتر از عملکرد الگوریتم، کل فرایند طراحی، داده، نظارت و تصمیم‌گیری را در بر گیرد.

از منظر نظریه خلق خطر، مسئولیت جبران خسارت متوجه کسی است که آن خطر را ایجاد کرده است. البته باید توجه داشت که خلق خطر، به تنهایی برای اثبات مسئولیت کیفری کافی نیست، اما می‌تواند مبنایی برای مسئولیت مدنی و نیز مسئولیت سازمانی فراهم آورد (هاشمی شاهرودی و میرزاامرجی: ۱۴۰۳).

۲-۴. کنترل مؤثر و مسئله کنترل انسانی معنادار

مفهوم «کنترل انسانی معنادار» که در مقدمه به آن اشاره شد، در اینجا جایگاه خود را در چارچوب نظری می‌یابد. ایده بنیادین این مفهوم آن است که برای اینکه یک تصمیم نظامی را بتوان از نظر اخلاقی و حقوقی به انسان نسبت داد، لازم است انسان «کنترل مؤثر» بر فرایند تصمیم‌گیری و اجرا داشته باشد (Horowitz, 2024: 310-313).

در سال‌های اخیر، مفهوم «کنترل انسانی معنادار» به یکی از محورهای اصلی مذاکرات بین‌المللی درباره سامانه‌های تسلیحاتی خودمختار تبدیل شده است. گروه کارشناسان دولتی کنوانسیون برخی سلاح‌های متعارف و کمیته بین‌المللی صلیب سرخ بر این نکته تأکید کرده‌اند که تصمیم درباره استفاده از زور باید همواره تحت نظارت و قضاوت واقعی انسان باقی بماند و اتکای صرف به خروجی سامانه‌های هوشمند، با الزامات حقوق بین‌الملل بشردوستانه سازگار نیست (ICRC, 2024؛ United Nations CCW-GGE, 2024).

کنترل مؤثر مستلزم سه شرط اصلی است. نخست اینکه انسان باید اطلاعات کافی و به‌هنگام از وضعیت داشته باشد. دوم اینکه باید فرصت و توانایی کافی برای ارزیابی مستقل آن اطلاعات و تصمیم‌گیری جایگزین داشته باشد. سوم اینکه باید بتواند در صورت لزوم، اجرای تصمیم را متوقف یا تغییر دهد (International Committee of the Red Cross (ICRC), 2021: 23-25). اگر اپراتور فرصت کافی برای ارزیابی مستقل اطلاعات یا امکان مداخله مؤثر در تصمیم را نداشته باشد، نمی‌توان از تحقق «کنترل انسانی معنادار» سخن گفت. در چنین وضعیتی، انسان بیش از آنکه «تصمیم‌گیرنده» باشد، «تأییدکننده خودکار» است و مسئولیت واقعی نمی‌تواند متوجه او گردد.

در همین راستا، هرچند در بسیاری از عملیات‌های مبتنی بر هوش مصنوعی، انسان به ظاهر در چرخه تصمیم‌گیری باقی می‌ماند، اما سرعت بسیار بالای پردازش داده‌ها، حجم گسترده اطلاعات و فشار فرماندهان برای شناسایی سریع اهداف، فرصت ارزیابی مستقل را از اپراتور سلب می‌کند. در چنین وضعیتی، حضور انسان بیش از آنکه بیانگر کنترل واقعی باشد، به تأییدی صوری بر تصمیم‌های الگوریتم تبدیل می‌شود و مفهوم «کنترل انسانی معنادار» با چالش جدی مواجه خواهد شد (Rona, 2026: 3).

در اینجا پرسش مهمی ظهور می‌کند: آیا فقدان کنترل انسانی معنادار به این معناست که دیگر هیچ انسانی مسئول نیست؟ پاسخ منفی است. اگر نشان دهیم که در حادثه‌ای مانند میناب، شرایط کنترل مؤثر برای اپراتور یا فرمانده فراهم نبوده، این مسئولیت را به طبقه‌های بالاتری از زنجیره تصمیم یعنی طراحان، سیاست‌گذاران و فرماندهان ارشد منتقل می‌کند. آن‌ها کسانی بودند که با طراحی فرایند عملیاتی به گونه‌ای که کنترل معنادار را از اپراتور سلب می‌کند، در واقع عاملیت انسانی را تضعیف کرده و زمینه را برای «واگذاری مسئولیت به ماشین» فراهم ساخته‌اند. از این رو، مفهوم کنترل انسانی معنادار نه راهی برای تبرئه انسان‌ها، بلکه ابزاری برای بازتوزیع مسئولیت به جایگاه‌ها و مقام‌های بالادستی است.

۲-۵. مفهوم مسئولیت شبکه‌ای در زنجیره تصمیم

پس از تبیین مفاهیم پیشین، اکنون می‌توان «مسئولیت شبکه‌ای» را به عنوان حلقه نهایی چارچوب نظری طرح کرد. این مفهوم در برابر الگوی سنتی مسئولیت فردی قرار می‌گیرد. در رویکرد کلاسیک، فرض بر آن است که باید «یک مقصر اصلی» شناسایی شود و دیگران، جز در موارد استثنایی، از دایره مسئولیت خارج می‌گردند. در مقابل، مسئولیت شبکه‌ای بر این پیش‌فرض استوار است که در سامانه‌های پیچیده‌ای مانند سامانه‌های تسلیحاتی مبتنی بر هوش مصنوعی، مسئولیت ویژگی یک شبکه از کنشگران به هم پیوسته است، نه صفت یک فرد منفرد (Fisher, 2020: 3120-3125).

در این شبکه، هر کنشگر، اعم از طراح الگوریتم، تأمین‌کننده داده‌های آموزشی، برنامه‌نویس رابط کاربری، اپراتور میدانی، فرمانده گردان، فرمانده کل و حتی نهاد ناظر و قانون‌گذار، سهمی در بروز خطا و نیز سهمی در پیشگیری از آن دارد. ارزیابی مسئولیت اخلاقی و حقوقی هر کنشگر باید متناسب با نقش و میزان تأثیر او در ایجاد یا افزایش خطر انجام شود.

با این نگاه، پرسش محوری دیگر این نیست که «چه کسی مقصر اصلی است؟» بلکه این است که «هر یک از کنشگران چه نسبتی از مسئولیت را بر عهده دارند؟» بنابراین در مورد حادثه میناب، در فرض احراز استفاده از سامانه‌های مبتنی بر هوش مصنوعی، طراح سامانه ممکن است از حیث

طراحی و تنظیم سامانه در معرض مسئولیت قرار گیرد. تأمین‌کننده داده‌های آموزشی ممکن است در صورت نقص یا سوگیری داده‌ها واجد مسئولیت باشد. برنامه‌نویس رابط کاربری ممکن است در صورت طراحی نامناسب سازوکارهای هشدار یا نظارت، در ارزیابی مسئولیت مورد توجه قرار گیرد. اپراتور ممکن است در حدود اختیارات و میزان کنترل مؤثر خود مسئول شناخته شود. فرماندهان نیز بسته به سطح تصمیم‌گیری، نحوه استقرار سامانه و رعایت الزامات احتیاطی، ممکن است واجد مسئولیت باشند.

در این برداشت، مسئولیت شبکه‌ای نه به معنای تبرئه همه است و نه به معنای سرزنش همه به یک میزان. هدف آن است که برای پر کردن شکاف مسئولیت، نظام بازخواست و پاسخگویی به گونه‌ای طراحی شود که تمامی کنشگرانی را که دارای «نقش علیّ معنادار» در شکل‌گیری ریسک و تحقق حادثه‌اند، متناسب با سهم واقعی‌شان در معرض پرسش و الزام قرار دهد. همچنین از منظر سیاست‌گذاری، این نگاه نشان می‌دهد که راهکار کاهش خطا نه در بهبود صرف الگوریتم، بلکه در بازطراحی کل شبکه تصمیم؛ از داده‌ها تا دستورالعمل‌های فرماندهی، نهفته است.

همچنین در تحلیل حقوقی این حادثه تصریح شده است که حتی اگر رفتار ارتكابی در همه فروض واجد شرایط تحقق جنایت جنگی نباشد، عدم رعایت اصل احتیاط در حمله همچنان می‌تواند موجب مسئولیت بین‌المللی دولت و در مواردی مسئولیت اشخاص تصمیم‌گیر شود. از این رو، ارزیابی حقوقی عملیات‌های مبتنی بر هوش مصنوعی نباید صرفاً بر وجود یا عدم وجود خطای الگوریتم متمرکز باشد، بلکه رعایت تکالیف احتیاطی و سازوکارهای نظارت انسانی نیز باید مورد توجه قرار گیرد (Rona, 2026: 4-6).

بدین ترتیب، با تکیه بر پنج مفهوم شکاف مسئولیت، علیت شبکه‌ای، نظریه خلق خطر، کنترل انسانی معنادار و مسئولیت شبکه‌ای، می‌توان نشان داد که فروکاستن تحلیل مسئولیت به «خطای هوش مصنوعی» تصویر کاملی از فرایند انتساب مسئولیت ارائه نمی‌کند و تحلیل شبکه‌ای ظرفیت بیشتری برای تبیین نقش کنشگران مختلف دارد.

۳. مطالعه موردی حادثه میناب

حادثه مدرسه شجره طیبه در میناب که در اسفند ۱۴۰۴ رخ داد، به دلیل پیامدهای گسترده انسانی، حقوقی و رسانه‌ای، یکی از مهم‌ترین نمونه‌های قابل بررسی در حوزه حقوق بین‌الملل بشردوستانه به‌شمار می‌رود. اگرچه در برخی گزارش‌های رسانه‌ای و تحلیل‌های کارشناسی احتمال نقش سامانه‌های مبتنی بر هوش مصنوعی در فرآیند شناسایی یا هدف‌یابی مطرح شده است، تاکنون هیچ گزارش رسمی و قابل‌آزمایی، استفاده از چنین سامانه‌هایی را در این عملیات به طور قطعی

تأیید نکرده است. از این رو، مطالعه حاضر می‌کوشد نشان دهد در فرض استفاده از سامانه‌های هوشمند، چارچوب مسئولیت حقوقی چگونه باید تحلیل شود.

۳-۱. شرح حادثه

در ۹ اسفند ۱۴۰۴ (۲۸ فوریه ۲۰۲۵)، در جریان موج نخست حملات هوایی ایالات متحده و اسرائیل علیه ایران، ساختمان دبستان و پیش‌دبستانی غیرانتفاعی «شجره طیبه»، واقع در محله مسکونی چمران، هدف اصابت موشک قرار گرفت. این مدرسه دو طبقه (همکف پسرانه و طبقه بالا دخترانه و مهدکودک) در آن لحظه میزبان بیش از ۱۵۰ دانش‌آموز چهار تا دوازده ساله، ۲۶ معلم زن و تعدادی از والدین بود.^۱ در این حمله، ۱۵۶ نفر جان باختند و ۹۵ نفر زخمی شدند.

شهر میناب در استان هرمزگان واقع شده و این استان به دلیل اشراف بر تنگه هرمز و خلیج فارس از اهمیت نظامی قابل توجهی برخوردار است؛ زیرا بخشی از فعالیت‌های اصلی نیروهای دریایی ایران در این منطقه متمرکز است. به گزارش شبکه الجزیره، مدرسه شجره طیبه میناب در زمره مدارسی بوده که وابستگی اداری به نیروی دریایی ایران داشته است؛ با این حال، این مدرسه صرفاً مختص فرزندان نیروهای دریایی نبوده است. همچنین بر اساس گزارش رسانه‌هایی چون رویترز، الجزیره و بی‌بی‌سی فارسی، این مدرسه پیش از سال ۱۳۹۵ در محدوده مجتمع نظامی سیدالشهدا قرار داشت، اما پس از آن از این مجموعه جدا و به صورت مستقل فعالیت می‌کرد.^۲ بنابراین، از منظر حقوقی، مدرسه شجره طیبه در مجاورت مقر تیپ ۳ نیروی دریایی قرار داشت و اصابت موشک کروز تاماهاوک به این ساختمان غیرنظامی می‌تواند پرسش‌هایی جدی در خصوص رعایت مفاد کنوانسیون‌های ژنو ۱۹۴۹، به ویژه اصل تمایز میان اهداف نظامی و غیرنظامی، مطرح سازد. بیانیه رسمی پنتاگون در ۱۰ اسفند ۱۴۰۴ علت حادثه را «انحراف فنی موشک» اعلام کرد، اما تحلیل برخی کارشناسان مستقل احتمال بروز خطای هدف‌گیری یا عدم رعایت کامل اصول تناسب و احتیاط را نیز مطرح کرده‌اند.^۳

این پژوهش روایت رسمی منتشرشده از سوی وزارت دفاع ایالات متحده مبنی بر «انحراف فنی موشک» را مبنای توصیف واقعه قرار می‌دهد، با این حال، از منظر حقوقی، این توضیح را پایان

۱. خبرگزاری تسنیم (۱۴۰۴)، «شهادت ۱۴۸ دانش‌آموز در حمله اسرائیل به میناب».

<https://www.tasnimnews.ir/fa/news/1404/12/D9%85%DB%8C%D9%86%D8%A7%D8%A8>
2. «Al Jazeera investigation: Iran girls' school targeting likely 'deliberate'»، Al Jazeera Media Network.

۳. درباره انگیزه این حمله، گمانه‌زنی‌های دیگری نیز مطرح شده است. واحد تحقیقات خبرگزاری الجزیره حمله به مدرسه را عمدی خوانده و دلیل آن را ایجاد شوک در میان مردم و تضعیف حمایت مردم ایران از نیروهای نظامی دانسته است.

تحلیل مسئولیت تلقی نمی‌کند. از این رو، صرف‌نظر از علت دقیق حادثه، ارزیابی حقوقی آن مستلزم بررسی نحوه تصمیم‌گیری، کیفیت اطلاعات، رعایت اصول حقوق بین‌الملل بشردوستانه و حدود مسئولیت کنشگران انسانی است.

۳-۲. پاسخ به روایت «خطای هوش مصنوعی»

پس از وقوع حادثه، در کنار روایت رسمی مبنی بر «انحراف فنی موشک»، برخی تحلیل‌های رسانه‌ای و حقوقی احتمال نقش سامانه‌های هوشمند در فرایند شناسایی یا هدف‌یابی را نیز مطرح کردند. در این قسمت، صرفاً در فرض استفاده از چنین سامانه‌هایی، پیامدهای حقوقی و آثار آن بر انتساب مسئولیت بررسی می‌شود.

روایت «خطای هوش مصنوعی»، حتی اگر صرفاً در سطح یک احتمال یا فرضیه مطرح شود، می‌تواند پیامدهای مهمی در تحلیل مسئولیت داشته باشد. تمرکز بر عملکرد الگوریتم، این خطر را ایجاد می‌کند که نقش تصمیم‌گیران انسانی، طراحان سامانه، فرماندهان نظامی و مقامات سیاسی در فرآیند تصمیم‌گیری کم‌رنگ شود و مسئولیت به‌جای توزیع در زنجیره تصمیم، به فناوری نسبت داده شود. در تحلیل حقوقی حادثه میناب تصریح شده است که حتی در فرض نقش داشتن سامانه‌های هوشمند، ارزیابی حقوقی باید همچنان بر تصمیم‌های انسانی، کیفیت اطلاعات، رعایت اصل احتیاط و میزان نظارت انسانی متمرکز باقی بماند و صرف اشاره به فناوری نمی‌تواند جایگزین بررسی مسئولیت اشخاص شود (Rona, 2026: 3-4).

همین رویکرد در اسناد بین‌المللی نیز پذیرفته شده است. کمیته بین‌المللی صلیب سرخ و گروه کارشناسان دولتی کنوانسیون برخی سلاح‌های متعارف بر این نکته تأکید کرده‌اند که استفاده از سامانه‌های هوشمند، پاسخگویی انسانی را از میان نمی‌برد و تصمیم درباره کاربرد زور باید همچنان تحت کنترل انسانی معنادار باقی بماند (ICRC, 2021: 23-25 ; United Nations CCW-). (GGE, 2024: 8-11).

این نتیجه از دو منظر قابل تأیید است. نخست، عملکرد سامانه‌های هوشمند متأثر از داده‌های آموزشی، طراحی الگوریتم، تنظیمات عملیاتی و نظارت انسانی است. دوم، حتی در فرض بروز خطا، مسئولیت ناشی از رعایت اصول تمایز، تناسب و احتیاط همچنان متوجه تصمیم‌گیران انسانی باقی می‌ماند (Rona, 2026: 4-5).

۳-۳. نقش سامانه‌های هوشمند

با گسترش استفاده از سامانه‌های مبتنی بر هوش مصنوعی در عملیات‌های نظامی، نقش این فناوری‌ها در چرخه شناسایی، تحلیل داده و انتخاب هدف به یکی از موضوعات مهم ادبیات

حقوق بین‌الملل بشردوستانه تبدیل شده است. گزارش‌های منتشرشده درباره برخی عملیات‌های نظامی ایالات متحده نشان می‌دهد که ارتش این کشور در سال‌های اخیر از سامانه‌های مبتنی بر یادگیری ماشین و تحلیل کلان‌داده برای پردازش تصاویر ماهواره‌ای، داده‌های پهبادی و اولویت‌بندی اهداف استفاده کرده است (The Guardian, 2025).

در فرض استفاده از سامانه‌های هوشمند در عملیات میناب، بررسی نقش این سامانه‌ها مستلزم توجه هم‌زمان به کیفیت داده‌های ورودی، نحوه طراحی الگوریتم، میزان نظارت انسانی و فرایند تصمیم‌گیری عملیاتی است.

در همین راستا، در تحلیل حقوقی این حادثه تأکید می‌شود که حتی اگر سامانه‌های مبتنی بر هوش مصنوعی در فرایند انتخاب هدف مورد استفاده قرار گرفته باشند، پرسش اصلی حقوقی نه عملکرد مستقل ماشین، بلکه نحوه تصمیم‌گیری انسان‌هایی است که داده‌ها را تولید، الگوریتم را طراحی، سامانه را به کارگیری و در نهایت حمله را تأیید کرده‌اند. به بیان دیگر، مسئولیت حقوقی همچنان باید بر زنجیره تصمیم‌گیری انسانی متمرکز باشد و نه بر فناوری به‌عنوان یک عامل مستقل (Rona, 2026: 3-4).

۴. بازشناسی مسئولیت در زنجیره تصمیم بر پایه چارچوب نظری پژوهش

پس از بررسی حادثه مدرسه شجره طیبه میناب، اکنون می‌توان به پرسش اصلی این مقاله بازگشت: چرا روایت «خطای هوش مصنوعی» نمی‌تواند پایان تحلیل مسئولیت باشد و در فرض استفاده از سامانه‌های هوشمند، مسئولیت چگونه باید در زنجیره تصمیم‌بازشناسی شود؟ برای پاسخ به این پرسش، ابتدا چارچوب نظری مقاله بر مطالعه موردی میناب تطبیق داده می‌شود، سپس نقش کنشگران انسانی در زنجیره تصمیم‌بررسی و در پایان، حادثه در پرتو اصول حقوق بین‌الملل بشردوستانه تحلیل می‌شود.

۴-۱. کاربست چارچوب نظری در حادثه میناب

چارچوب نظری مسئولیت در سامانه‌های هوش مصنوعی نظامی که در بخش دوم این مقاله تدوین شد، امکان تحلیل عمیق‌تر حادثه میناب را فراهم می‌کند.

نخست، در فرض استفاده از سامانه‌های هوشمند در فرایند هدف‌یابی، مسئله «شکاف مسئولیت» به‌خوبی قابل تصور است. در چنین فرضی، تصمیم‌نهایی محصول تعامل مجموعه‌ای از کنشگران انسانی و فرایندهای فنی است و به دشواری می‌توان یک «فاعل انسانی واحد» را مسئول تمامی پیامدهای عملیات دانست. همین وضعیت، زمینه شکل‌گیری «شکاف مسئولیت» را فراهم می‌کند و ضرورت تحلیل مسئولیت در قالب یک شبکه تصمیم را نشان می‌دهد.

دوم، مفهوم «علیت شبکه‌ای» در تحلیل حادثه میناب زمانی معنا پیدا می‌کند که استفاده از سامانه‌های هوشمند در فرایند شناسایی یا انتخاب هدف مفروض گرفته شود. در چنین سناریویی، حادثه حاصل برهم کنش مجموعه‌ای از عوامل، از جمله کیفیت و به‌روز بودن داده‌های اطلاعاتی، طراحی و تنظیم الگوریتم‌ها، سطح نظارت انسانی، فشار زمانی بر اپراتورها و ساختار تصمیم‌گیری نظامی خواهد بود. از این‌رو، تمرکز صرف بر «خطای الگوریتم» تصویری ناقص از زنجیره علّی حادثه ارائه می‌کند.

سوم، نظریه «خلق خطر» نیز دولت استفاده‌کننده از سامانه را در معرض مسئولیت قرار می‌دهد. دولت مهاجم با تصمیم آگاهانه برای استفاده از سامانه‌های هوش مصنوعی در عملیات نظامی، آن هم در محیطی که حضور گسترده غیرنظامیان در آن قابل پیش‌بینی بوده است، خطری فراتر از حد متعارف برای جمعیت غیرنظامی ایجاد کرده است (پیری، ۱۴۰۳: ۱۳). این خلق خطر، حتی مستقل از شناسایی یک عامل انسانی مشخص به‌عنوان مباشر، می‌تواند از منظر حقوقی در چارچوب تعهدات مربوط به رعایت اصل احتیاط در حمله تحلیل شود.

چهارم، در فرض استفاده از سامانه‌های هوشمند، پرسش اساسی آن است که آیا تصمیم‌گیرندگان انسانی فرصت و امکان ارزیابی مستقل خروجی سامانه را داشته‌اند یا خیر. از این منظر، مسئله اصلی نه حذف انسان از چرخه تصمیم، بلکه تضعیف «کنترل انسانی معنادار» بر تصمیم‌نهایی است (Rona, 2026).

در نهایت، مفهوم «مسئولیت شبکه‌ای» چارچوب مناسب‌تری برای تحلیل حادثه میناب فراهم می‌کند. در این رویکرد، به‌جای جست‌وجوی یک مقصر واحد، مسئولیت بر اساس نقش هر یک از کنشگران انسانی در زنجیره تصمیم ارزیابی می‌شود. این رویکرد می‌تواند از بی‌کیفیری جلوگیری کرده و پاسخگویی را در عملیات‌های مبتنی بر سامانه‌های هوشمند تقویت کند.

۴-۲. تحلیل مسئولیت انسانی

با تکیه بر چارچوب نظری فوق، اکنون می‌توان در فرض استفاده از سامانه‌های هوشمند در فرآیند هدف‌یابی، نقش چهار گروه اصلی کنشگران انسانی را در زنجیره علّی حادثه میناب بررسی کرد. هدف این تحلیل، انتساب قطعی مسئولیت به اشخاص یا نهادهای مشخص نیست، بلکه نشان دادن آن است که حتی در صورت به‌کارگیری سامانه‌های مبتنی بر هوش مصنوعی، مسئولیت حقوقی همچنان در میان کنشگران انسانی توزیع می‌شود و به فناوری منتقل نمی‌گردد (Rona, 2026: 4-5).

در وهله نخست، طراحان و توسعه‌دهندگان سامانه در صورتی که سامانه‌های هوشمند در عملیات به‌کار گرفته شده باشند، در مرحله طراحی الگوریتم، تعیین معیارهای طبقه‌بندی اهداف،

تنظیم آستانه‌های تصمیم‌گیری و پیش‌بینی سازوکارهای بازیابی انسانی نقش اساسی دارند. در ادبیات حقوقی نیز تأکید شده است که طراحی سامانه‌های پرخطر باید به گونه‌ای باشد که احتمال آسیب به غیرنظامیان تا حد امکان کاهش یابد و امکان مداخله مؤثر انسان در تصمیم‌نهایی حفظ شود (Rona, 2026: 3-4).

در مرحله بعد، تولیدکنندگان و مدیران داده‌های آموزشی و اطلاعات عملیاتی نیز از ارکان مهم زنجیره تصمیم‌به‌شمار می‌آیند. چنانچه سامانه‌ای بر پایه داده‌های ناقص، قدیمی یا نادرست آموزش دیده یا تغذیه شده باشد، احتمال بروز خطا در شناسایی هدف افزایش می‌یابد. در تحلیل‌های منتشرشده درباره حادثه میناب نیز این احتمال مطرح شده که در صورت استفاده از سامانه‌های هوشمند، کیفیت داده‌های اطلاعاتی و به‌روزرسانی مستمر آن‌ها می‌توانسته بر نتیجه عملیات تأثیرگذار باشد؛ هرچند این موضوع تاکنون به‌صورت رسمی تأیید نشده است (Rona, 2026: 2-3).

اپراتور انسانی نیز، حتی در پیشرفته‌ترین سامانه‌های هوشمند، همچنان یکی از حلقه‌های اصلی مسئولیت محسوب می‌شود. مطابق حقوق بین‌الملل بشردوستانه، اتکای صرف به خروجی سامانه جایگزین تکلیف انسانی برای ارزیابی مشروعیت هدف نیست. از این‌رو، در فرض استفاده از سامانه‌های هوشمند، مسئولیت اپراتور وابسته به میزان اطلاعات در اختیار، فرصت ارزیابی مستقل و امکان اعمال کنترل مؤثر بر تصمیم‌نهایی خواهد بود (ICRC, 2021: 23-25).

در نهایت، فرماندهان نظامی و مقامات تصمیم‌گیر گسترده‌ترین سطح مسئولیت را بر عهده دارند؛ زیرا تصمیم به استقرار سامانه، تعیین قواعد استفاده، ایجاد سازوکارهای نظارتی و اطمینان از رعایت اصول حقوق بین‌الملل بشردوستانه در صلاحیت آنان قرار دارد. همان‌گونه که کمیته بین‌المللی صلیب سرخ و تحلیل‌های حقوقی اخیر درباره حادثه میناب تأکید کرده‌اند، حتی اگر در یک عملیات از سامانه‌های هوشمند استفاده شود، این امر مسئولیت فرماندهان را منتفی نمی‌کند و آنان همچنان مکلف به تضمین رعایت اصول تفکیک، تناسب و احتیاط هستند (ICRC, 2021: 25-27).

بر این اساس، در فرض استفاده از سامانه‌های هوشمند، ارزیابی مسئولیت مستلزم بررسی نقش هر یک از کنشگران انسانی در زنجیره تصمیم، کیفیت داده‌های مورد استفاده، نحوه طراحی سامانه و میزان نظارت انسانی است. در این چارچوب، رویکرد «مسئولیت شبکه‌ای» امکان می‌دهد مسئولیت حقوقی و اخلاقی هر یک از کنشگران متناسب با نقش واقعی آنان در فرآیند تصمیم‌گیری ارزیابی شود، بی‌آنکه مسئولیت به فناوری به‌عنوان یک عامل مستقل منتقل گردد.

۴-۳. تحلیل حقوقی در چارچوب حقوق بین‌الملل بشردوستانه

در راستای بازشناسی مسئولیت در زنجیره تصمیم، حادثه میناب در پرتو چهار اصل بنیادین حقوق بین‌الملل بشردوستانه، یعنی تفکیک، تناسب، احتیاط در حمله و مسئولیت فرماندهی، مورد ارزیابی قرار می‌گیرد.

۴-۳-۱. اصل تفکیک

اصل تفکیک بی‌تردید یکی از بنیادی‌ترین اصول حقوق بین‌الملل بشردوستانه به‌شمار می‌رود. بر اساس این اصل، طرف‌های درگیر در یک مخاصمه مسلحانه موظف هستند همواره میان افراد غیرنظامی و رزمندگان و نیز میان اهداف نظامی و اموال غیرنظامی تمایز قائل شوند (اعلانی‌فرد، انصاری مهبیاری و راعی دهقی، ۱۴۰۳: ۸-۱۰). در نتیجه، حمله به افراد یا اموال غیرنظامی به‌عنوان هدف اصلی به‌طور کلی ممنوع است.

مهم‌ترین سندی که این اصل را به‌صراحت بیان کرده، پروتکل الحاقی اول به کنوانسیون‌های ژنو مصوب ۱۹۷۷ است. ماده ۴۸ این پروتکل اصل تفکیک را به‌عنوان «قاعده بنیادین» معرفی می‌کند و مقرر می‌دارد که طرف‌های مخاصمه باید همواره میان جمعیت غیرنظامی و رزمندگان و نیز میان اهداف نظامی و اموال غیرنظامی تمایز قائل شوند. افزون بر این، بند ۲ ماده ۵۱ همان پروتکل تصریح می‌کند که «جمعیت غیرنظامی و افراد غیرنظامی به‌طور کلی نباید مورد حمله قرار گیرند» و هرگونه اعمال خشونت‌آمیز یا تهدیدهایی که هدف اصلی آن ایجاد رعب و وحشت در میان جمعیت غیرنظامی باشد ممنوع است.^۱ این اصل همچنین در کنوانسیون‌های چهارگانه ژنو ۱۹۴۹ و در حقوق بین‌الملل عرفی نیز به رسمیت شناخته شده است (ICRC 2005:3-24).

برای تشخیص اینکه آیا حمله به یک مکان یا شیء نقض اصل تفکیک محسوب می‌شود یا خیر، نخست باید مفهوم «هدف نظامی» روشن شود. ماده ۵۲ پروتکل الحاقی اول هدف نظامی را چنین تعریف می‌کند: اشیایی که بر اساس ماهیت، موقعیت، هدف یا کاربرد خود به‌طور مؤثر در اقدامات نظامی مشارکت دارند و انهدام کامل یا جزئی، تصرف یا خنثی‌سازی آن‌ها مزیت نظامی مشخصی فراهم می‌آورد.^۲ بر این اساس، تأسیسات نظامی، تجهیزات جنگی، مراکز فرماندهی، پادگان‌ها و نیروهای رزمی در زمره اهداف نظامی قرار می‌گیرند. در مقابل، اشیایی که چنین نقشی در عملیات نظامی ندارند، از حمایت برخوردارند و نباید هدف حمله قرار گیرند.

مدارس به‌دلیل ماهیت آموزشی و حضور کودکان در آن‌ها معمولاً در شمار اماکن غیرنظامی

1. ICRC (1977), Protocol Additional to the Geneva Conventions of 12 August 1949, and relating to the Protection of Victims of International Armed Conflicts (Protocol I), Article 51(2). Available at: <https://ihl-databases.icrc.org/en/ihl-treaties/api-1977/article-51>

2. ICRC (1977), Protocol I, Article 52(2).

قرار می‌گیرند و از حمایت ویژه برخوردارند. هرچند کنوانسیون‌های ژنو و پروتکل‌های الحاقی برای مدارس مصونیت مطلق مقرر نکرده‌اند،^۱ اما در عمل این اماکن در زمره فضاهایی محسوب می‌شوند که باید با نهایت احتیاط با آن‌ها برخورد شود. افزون بر این، اسنادی همچون کنوانسیون حقوق کودک ۱۹۸۹ و پروتکل اختیاری آن درباره مشارکت کودکان در مخاصمات مسلحانه بر ضرورت حمایت ویژه از کودکان در زمان جنگ تأکید کرده‌اند.^۲

با توجه به اطلاعات منتشرشده درباره حادثه میناب، این واقعه از منظر اصل تفکیک شایسته بررسی دقیق حقوقی است. بر اساس گزارش‌های منتشرشده، مدرسه در زمان حمله یک مرکز آموزشی فعال بوده و کودکان و کارکنان آموزشی در آن حضور داشته‌اند. همچنین تاکنون هیچ گزارش رسمی و قابل راستی‌آزمایی مبنی بر استفاده نظامی از ساختمان مدرسه در زمان حمله منتشر نشده است؛ هرچند برخی منابع به سابقه استقرار این مدرسه در مجاورت یک مجموعه نظامی اشاره کرده‌اند (Rona, 2026: 3-4؛ ICRC, 2020: 18-20).

در چنین فرضی، صرف سابقه استقرار مدرسه در محدوده یک مجموعه نظامی یا مجاورت آن با تأسیسات نظامی، به تنهایی موجب از بین رفتن حمایت حقوقی آن نمی‌شود؛ زیرا مطابق ماده ۵۲ پروتکل الحاقی اول، ملاک، کاربری مؤثر و بالفعل هدف در زمان حمله است، نه پیشینه تاریخی یا موقعیت جغرافیایی آن (Protocol I, Art.52؛ ICRC, 2005: 25-27). از این رو، اگر احراز شود که ساختمان در زمان حمله کاربری آموزشی داشته و در عملیات نظامی مشارکت مؤثری نداشته است، حمله به آن می‌تواند از منظر اصل تفکیک محل تردید جدی باشد.

۴-۳-۲. اصل تناسب

اصل تناسب یکی از پیچیده‌ترین اصول حقوق بین‌الملل بشردوستانه به‌شمار می‌رود. این اصل ناظر بر حملاتی است که علیه اهداف نظامی مشروع انجام می‌شوند اما در عین حال ممکن است موجب تلفات یا خسارات جانبی برای غیرنظامیان شوند. پرسش اساسی در اینجا آن است که آیا مزیت نظامی مورد انتظار از حمله به اندازه‌ای قابل توجه است که خسارات جانبی احتمالی را توجیه کند یا خیر.

مبنای حقوقی این اصل در بند ۵ (ب) ماده ۵۱ پروتکل الحاقی اول ۱۹۷۷ بیان شده است. مطابق این مقرر، حملاتی که انتظار می‌رود به طور اتفاقی موجب تلفات جانی غیرنظامیان،

۱. چراکه اگر مدرسه واقعاً توسط نظامیان به عنوان پناهگاه یا مرکز فرماندهی استفاده شود، ممکن است هدف نظامی تلقی گردد.

2. Optional Protocol to the Convention on the Rights of the Child on the involvement of children in armed conflict (2000), Articles 1-4.

مجروح شدن آنان یا آسیب به اموال غیرنظامی شود، در صورتی که این خسارات در مقایسه با مزیت نظامی مشخص و مستقیم مورد انتظار زیاده‌روانه باشد، ممنوع است. این اصل در ماده ۵۷ همان پروتکل نیز مورد تأکید قرار گرفته است. همچنین مطالعه عرفی کمیته بین‌المللی صلیب سرخ، اصل تناسب را یکی از قواعد مسلم حقوق بین‌الملل بشردوستانه دانسته که رعایت آن در تمامی مخاصمات مسلحانه الزامی است (ICRC, 2005: 46-49).

در خصوص حادثه میناب، تحلیل اصل تناسب وابسته به احراز ماهیت هدف مورد حمله است. اگر مطابق گزارش‌های منتشرشده، ساختمان مدرسه در زمان حمله صرفاً کاربری آموزشی داشته و هدف نظامی محسوب نمی‌شده، بحث اصلی در وهله نخست به اصل تفکیک مربوط خواهد بود. اما چنانچه فرض شود تصمیم‌گیرندگان نظامی، به هر دلیل، ساختمان را به عنوان یک هدف نظامی تلقی کرده‌اند، آنگاه ارزیابی تناسب اهمیت می‌یابد (Rona, 2026: 4-5).

در این فرض، باید بررسی شود که آیا مزیت نظامی مورد انتظار از حمله، در مقایسه با خطر قابل پیش‌بینی وارد آمدن خسارات به غیرنظامیان، به‌ویژه کودکان حاضر در مدرسه، متناسب بوده است یا خیر. با توجه به اطلاعات منتشرشده درباره شمار قربانیان و ماهیت محل اصابت، این موضوع می‌تواند تردیدهای جدی درباره رعایت اصل تناسب ایجاد کند؛ هرچند اظهارنظر قطعی در این خصوص مستلزم دسترسی به اطلاعات عملیاتی، ارزیابی اطلاعاتی فرماندهان و نتایج تحقیقات مستقل است (Rona, 2026: 5-6).

از این رو، حادثه میناب را می‌توان از دو منظر در ارتباط با اصل تناسب بررسی کرد: نخست، اگر احراز شود که هدف اساساً نظامی نبوده است، بحث اصلی ناظر بر اصل تفکیک خواهد بود؛ دوم، اگر فرض شود هدف بر مبنای اطلاعات موجود برای تصمیم‌گیرندگان نظامی، یک هدف نظامی تلقی شده، باید بررسی شود که آیا خسارات قابل پیش‌بینی واردشده به غیرنظامیان در مقایسه با مزیت نظامی مورد انتظار، زیاده‌روانه بوده است یا خیر. در هر دو فرض، استفاده احتمالی از سامانه‌های هوش مصنوعی تأثیری بر اصل مسئولیت انسانی در ارزیابی تناسب ندارد و تصمیم‌گیرندگان انسانی همچنان مکلف به رعایت این اصل هستند (ICRC, 2021: 23-27).

۴-۳-۳. اصل احتیاط

اصل احتیاط در حمله یکی از عملی‌ترین اصول حقوق بین‌الملل بشردوستانه است و بر تعهد فعال طرف‌های مخاصمه برای اتخاذ «همه اقدامات ممکن» جهت کاهش آسیب به غیرنظامیان تأکید دارد. مستند اصلی این اصل ماده ۵۷ پروتکل الحاقی اول ۱۹۷۷ با عنوان «احتیاط در حمله» است. بر اساس بند ۱ این ماده، در جریان عملیات نظامی باید همواره مراقبت‌های لازم برای حفظ

امنیت جمعیت غیرنظامی، افراد غیرنظامی و اموال غیرنظامی به عمل آید. همچنین بند ۲ (الف) (۱) مقرر می‌دارد که کسانی که عملیات نظامی را برنامه‌ریزی یا اجرا می‌کنند، باید همه اقدامات ممکن را برای اطمینان از این موضوع انجام دهند که هدف مورد نظر یک هدف نظامی مشروع است. بند ۲ (الف) (۲) نیز تأکید می‌کند که طرف‌های مخاصمه باید روش‌ها و ابزارهای حمله‌ای را انتخاب کنند که احتمال وارد آمدن تلفات به غیرنظامیان را به حداقل برساند. افزون بر این، بند ۲ (ب) تصریح می‌کند که هرگاه مشخص شود هدف غیرنظامی است یا درباره ماهیت آن تردید معقول وجود دارد، حمله باید لغو یا متوقف شود. این تعهد در مطالعه عرفی کمیته بین‌المللی صلیب سرخ نیز به عنوان یکی از قواعد بنیادین حقوق بین‌الملل بشردوستانه مورد تأکید قرار گرفته است (ICRC, 2005: 51-54).

اصطلاح «همه اقدامات ممکن» به اقداماتی اشاره دارد که با توجه به اطلاعات در دسترس، امکانات موجود و شرایط زمان تصمیم‌گیری، به‌طور منطقی از برنامه‌ریزان و مجریان عملیات انتظار می‌رود. این اقدامات شامل گردآوری و راستی‌آزمایی اطلاعات، ارزیابی مجدد هدف، انتخاب شیوه و وسیله مناسب حمله و در صورت لزوم لغو یا تعلیق عملیات است (ICRC, 2016: 203-206).

در خصوص حادثه میناب، چند پرسش حقوقی قابل طرح است. نخست، اگر اطلاعات مربوط به ماهیت هدف یا پایگاه‌های داده عملیاتی به‌روز نبوده باشد، باید بررسی شود که آیا تمامی اقدامات ممکن برای راستی‌آزمایی اطلاعات پیش از حمله انجام شده است یا خیر. دوم، در فرض استفاده از سامانه‌های هوشمند، باید ارزیابی شود که آیا اپراتورها و فرماندهان امکان واقعی برای بازبینی مستقل خروجی سامانه و توقف عملیات در صورت وجود تردید را داشته‌اند یا خیر.

سوم، ماده ۳۶ پروتکل الحاقی اول دولت‌ها را مکلف می‌کند پیش از به‌کارگیری سلاح‌ها یا روش‌های جدید جنگی، سازگاری آن‌ها را با حقوق بین‌الملل بشردوستانه ارزیابی کنند. از این رو، اگر در این عملیات از سامانه‌های مبتنی بر هوش مصنوعی استفاده شده باشد، این پرسش مطرح می‌شود که آیا ارزیابی حقوقی لازم درباره قابلیت اطمینان، دقت سامانه، کیفیت داده‌های مورد استفاده و امکان اعمال کنترل انسانی معنادار پیش از استقرار آن انجام شده است یا خیر (ICRC, 2021: 15-18).

در نهایت، اصل احتیاط اقتضا می‌کند که هرگاه درباره ماهیت غیرنظامی یا نظامی یک هدف تردید معقول وجود داشته باشد، عملیات متوقف یا دست‌کم مورد بازبینی مجدد قرار گیرد. بنابراین، اگر نشانه‌های قابل تشخیص از کاربری آموزشی ساختمان یا حضور غیرنظامیان وجود داشته باشد، رعایت اصل احتیاط محل ارزیابی حقوقی خواهد بود؛ هرچند نتیجه‌گیری قطعی مستلزم دسترسی به اسناد عملیاتی و تحقیقات مستقل است.

نتیجه‌گیری

بررسی حادثه مدرسه شجره طیبه میناب نشان داد که تقلیل این رویداد به «خطای هوش مصنوعی» نه تنها از نظر فنی تحلیلی ناقص است، بلکه از منظر حقوق بین‌الملل بشردوستانه و نظریه‌های نوین مسئولیت نیز نمی‌تواند مبنای مناسبی برای انتساب یا رفع مسئولیت باشد. سامانه‌های هوش مصنوعی نظامی در خلأ تصمیم‌گیری نمی‌کنند، بلکه بر پایه داده‌هایی آموزش می‌بینند که انسان‌ها گردآوری و اعتبارسنجی کرده‌اند، مطابق معیارهایی عمل می‌کنند که طراحان تعیین کرده‌اند و در قالب دستورالعمل‌هایی به کار گرفته می‌شوند که توسط فرماندهان و نهادهای نظامی تصویب شده است. از این رو، آنچه در ظاهر «خطای الگوریتم» تلقی می‌شود، در واقع محصول مجموعه‌ای از تصمیمات انسانی در مراحل طراحی، آموزش، استقرار و نظارت بر سامانه است.

کاربست چارچوب نظری پژوهش بر مطالعه موردی میناب نشان داد که مفاهیمی همچون «شکاف مسئولیت»، «علیت شبکه‌ای»، «خلق خطر»، «کنترل انسانی معنادار» و «مسئولیت شبکه‌ای» ظرفیت مناسبی برای تبیین مسئولیت در عملیات‌های نظامی مبتنی بر فناوری‌های نوین دارند. بر این اساس، در فرض استفاده از سامانه‌های هوشمند، مسئولیت حقوقی نباید به عملکرد سامانه یا یک عامل منفرد تقلیل یابد، بلکه باید نقش طراحان سامانه، تولیدکنندگان و مدیران داده، تحلیلگران اطلاعات، اپراتورها، فرماندهان نظامی و دولت استفاده‌کننده، متناسب با سهم هر یک در فرایند تصمیم‌گیری، مورد ارزیابی قرار گیرد.

تحلیل حادثه در پرتو اصول حقوق بین‌الملل بشردوستانه نیز نشان داد که توسعه فناوری‌های هوش مصنوعی تغییری در الزامات بنیادین این نظام حقوقی ایجاد نمی‌کند. اصول تفکیک، تناسب، احتیاط در حمله و مسئولیت فرماندهی همچنان معیارهای اصلی ارزیابی مشروعیت عملیات‌های نظامی باقی می‌مانند و بهره‌گیری احتمالی از سامانه‌های هوشمند، این تعهدات را از میان نمی‌برد. در نتیجه، پاسخگویی حقوقی همچنان متوجه انسان‌ها و دولت استفاده‌کننده از سامانه است و فناوری به‌تنهایی نمی‌تواند موضوع انتساب مسئولیت قرار گیرد.

یافته‌های این پژوهش همچنین نشان می‌دهد که چالش اصلی هوش مصنوعی در مخاصمات مسلحانه، فقدان قواعد حقوقی نیست، بلکه دشوارتر شدن فرایند انتساب مسئولیت در زنجیره‌های پیچیده تصمیم‌گیری است. از این رو، به جای ایجاد گسست در قواعد موجود، ضروری است مفاهیم حقوق بین‌الملل بشردوستانه با اقتضائات فناوری‌های نوین بازخوانی و بازتفسیر شوند؛ به‌ویژه مفهوم «کنترل انسانی معنادار» که می‌تواند معیار مهمی برای ارزیابی مشروعیت استفاده از سامانه‌های هوشمند و تعیین حدود مسئولیت کنشگران انسانی باشد.

بر این اساس، کاهش شکاف مسئولیت در سامانه‌های هوشمند نظامی مستلزم تقویت

ارزیابی‌های حقوقی پیش از استقرار این سامانه‌ها، تضمین حفظ کنترل انسانی معنادار در تمامی مراحل تصمیم‌گیری و ایجاد سازوکارهای شفاف پاسخگویی برای تمامی کنشگران دخیل در زنجیره تصمیم است. تنها با چنین رویکردی می‌توان از تبدیل روایت «خطای هوش مصنوعی» به ابزاری برای تضعیف پاسخگویی حقوقی جلوگیری کرد و هم‌زمان، کارآمدی قواعد حقوق بین‌الملل بشردوستانه را در مواجهه با فناوری‌های نوین حفظ نمود.

کتابنامه

- اعلانی‌فرد، سپیده؛ انصاری مهیاری، علیرضا؛ راعی دهقی، مسعود (۱۴۰۳). اصل تناسب در پرتو استفاده از هوش مصنوعی در مخاصمات مسلحانه. *مجله مطالعات حقوقی فضای مجازی*، ۳ (۱).
- انصاری مهیاری، علیرضا؛ محمودی، هادی (۱۴۰۱). بررسی راهکارهای تقلیل حملات سایبری از منظر حقوق بین‌الملل بشردوستانه. *مجله مطالعات حقوقی فضای مجازی*، ۱ (۳).
- بنافی، فرشته (۱۴۰۲). حفاظت از حق حریم خصوصی اطلاعاتی در مقابل تهدیدات ناشی از هوش مصنوعی نظامی. *مجله پژوهش حقوق خصوصی*، ۱۲ (۴۵).
- بیگلربیگی، کیان؛ عزیزی، ستار (۱۴۰۴). شرایط و الزامات حاکم بر عملیات سایبری در زمان اشغال نظامی در پرتو راهنمای تالین ۲، *مجله حقوق و مطالعات سیاسی*، ۵ (۱).
- خلیفه، اسماعیل؛ اسماعیل‌تبار، احمد (۱۴۰۲). اشاعه تسلیحات تهاجمی خودکار در حقوق بین‌الملل بشردوستانه اسلامی. *مجله پژوهش‌های فقهی*، ۱۹ (۲).
- زبده، محمد؛ حقیقت‌پور، حسین (۱۴۰۲). مطالعه تطبیقی در مصونی لآت غیرنظامیان و زوال آن در نظام‌های حقوق بشردوستانه اسلام و بین‌الملل. *مجله مطالعات حقوق بشر/اسلامی*، ۱۲ (۳).
- زمانی، سید قاسم؛ طلایی، فرهاد؛ بلاغی، ابوذر (۱۴۰۴). رویکرد حمایتی دیوان کیفری بین‌المللی از حقوق بشردوستانه کودکان در تخلفات شدید. *مجله پژوهش‌های حقوق جزا و جرم‌شناسی*، ۱۳ (۲۶).
- زمانیان، معصومه؛ وطنی، زهرا (۱۴۰۴). نقش سلاح‌های خودکار مبتنی بر هوش مصنوعی و کنترل انسانی در حقوق بین‌الملل بشردوستانه: چالش یا تعامل؟ *مجله مطالعات حقوقی*، ۱۷ (۴).
- سیدناصری، محمدمهدی؛ درویشی، سینا (۱۴۰۲). حمایت از حقوق کودکان در مخاصمات مسلحانه بین‌المللی از منظر حقوق بین‌الملل بشردوستانه. *مجله حقوق و مطالعات نوین*، ۱۱.
- صادق‌نیا، مجتبی؛ امجدیان، فرامرز (۱۴۰۴). تحلیل عملکرد بین‌المللی دولت‌ها در حوزه جنگ و منازعات. *مجله قانون‌یار*، ۹ (۳۳).

غلام‌حیدری، صالح (۱۴۰۳). امکان‌سنجی دفاع مشروع از طریق حملات سایبری. *مجله تمدن حقوقی*، ۷ (۲۰).

قلندری شمایی، طاهر؛ سبحانی، مهین (۱۴۰۴). سپر انسانی و چالش‌های پیش روی آن با نگاهی به اصل تناسب در حقوق بین‌الملل بشردوستانه با تأکید بر حمله سال ۲۰۲۳-۲۰۲۴ اسرائیل به غزه. *مجله حقوقی بین‌المللی*، ۴۲ (۷۸).

گلچین، پدرام؛ نجفی، فیروز؛ شهبازی، آرامش (۱۴۰۱). حقوق کودکان در عملیات انتحاری: مطالعه موردی بوکوحرام. *مجله حقوق پزشکی*، ۱۶.

متین، محمدسکندر؛ اطهر، حمیده (۱۴۰۳). تحلیل حقوقی مخاصمات مسلحانه سال ۲۰۲۳ اسرائیل و حماس (گروه مقاومت اسلامی و جهادی). *مجله دانش حقوقی*، ۲ (۶).

محمودی، هادی؛ انصاری مهباری، علیرضا؛ راعی دهقی، مسعود (۱۴۰۴). تحلیل نقش هوش مصنوعی در بهبود اجرای اصول بنیادین مخاصماتی. *مجله مطالعات حقوقی فضای مجازی*، ۴ (۲).

موثقی، حسن (۱۴۰۱). چالش‌های کاربرد حقوق بین‌الملل بشردوستانه در جنگ‌های سایبری. *مجله مطالعات حقوقی فضای مجازی*، ۱ (۲).

موسی‌زاده، رضا؛ آذرپندار، احمدرضا (۱۴۰۳). جنگ‌افزارهای خودفرمان از منظر حقوق بین‌الملل بشردوستانه. *مجله مطالعات حقوق عمومی*، ۵۴ (۳).

نوروزی، میثم؛ اسکندری خوشگو، مهدی؛ ابوالقاسمی، ساناز (۱۴۰۲). مطالعه تطبیقی رویکرد فقه اسلامی و حقوق بشردوستانه در حمایت از حقوق کودکان در جنگ. *مجله تمدن حقوقی*، ۶ (۱۴).

نوروزی، نیما؛ درویشی، اطهر؛ قاسم‌تبار، سیدمحمدرضا (۱۴۰۴). تحلیل تطبیقی نقض حقوق کودکان در مخاصمات مسلحانه و تطبیق آن با تروریسم دولتی: مطالعه موردی اقدامات رژیم صهیونیستی با تأکید بر حملات اخیر به ایران در پرتو حقوق بین‌الملل بشردوستانه. *مجله قضایانه*، ۳ (۸).

هاشمی شاهرودی، سیدمحمدصادق؛ میرزاامرجی، مهنوش (۱۴۰۳). مسئولیت دولت‌ها در قبال نقض حقوق بین‌المللی بشردوستانه غیرنظامیان در مخاصمات مسلحانه. *مجله مطالعات قدرت نرم*، ۱۴ (۳).

Aiden Warren, Alek Hillas (2018). LETHAL AUTONOMOUS ROBOTICS: RETHINKING THE DEHUMANIZATION OF WARFARE. *UCLA Journal of International Law and Foreign Affairs*, 22 (2).

Asaro, Peter M. (2022). The Liability Problem for Autonomous Artificial Agents. In: *The Oxford Handbook of AI Governance*, Oxford University Press.

- C. Anthony Pfaff (2019). The Ethics of Acquiring Disruptive Technologies: Artificial Intelligence, Autonomous Weapons, and Decision Support Systems. *PRISM*, 8 (3).
- Can Kasapoğlu, Barış Kırdemir (2019). *WARS OF NONE: ARTIFICIAL INTELLIGENCE AND THE FUTURE OF CONFLICT*. Centre for Economics and Foreign Policy Studies
- Christina Balis, Paul O'Neill (2022). *AI and Human Agency*. Trust in AI: Rethinking Future Command, Royal United Services Institute (RUSI).
- Crystal S. Yang, Will Dobbie (2020). EQUAL PROTECTION UNDER ALGORITHMS: A NEW STATISTICAL AND LEGAL FRAMEWORK. *Michigan Law Review*, 119 (2).
- Fisher, Elizabeth (2020). Reconstructing Responsibility in Automated Decision-Making Systems. *Modern Law Review*, 83 (6).
- H. Akin Ünver (2018). *Artificial Intelligence, Authoritarianism and the Future of Political Systems*. Centre for Economics and Foreign Policy Studies.
- Latour, Bruno (2005). *Reassembling the Social: An Introduction to Actor-Network-Theory*. Oxford: Oxford University Press.
- Lu, Sylvia (2022), *Data Privacy, Human Rights, and Algorithmic Opacity*, *California Law Review*, Vol. 110, No. 6, pp. 2087-2147
- Maria Stefania Cataleta (2020), *Humane Artificial Intelligence: The Fragility of Human Rights Facing AI*, East-West Center
- Mary E. Webb, Andrew Fluck, Johannes Magenheimer, Joyce Malyn-Smith, Juliet Waters, Michelle Deschênes, Jason Zagami (2021), *Machine learning for human learners: opportunities, issues, tensions and threats*, Educational Technology Research and Development, Vol. 69, No. 4, Special Issue: Embodied Cognition and Technology for Learning | Special Issue: Learners and Learning Contexts: International Perspectives on New Alignments for the Digital Age, pp. 2109-2130
- Nadia Marsan, Steven Hill (2019), *International law and military applications of Artificial Intelligence*, The Brain and the Processor: Unpacking the Challenges of Human-Machine Interaction, NATO Defense College
- NICO LÜCK (2019), *MACHINE LEARNING-POWERED ARTIFICIAL INTELLIGENCE IN WEAPONS SYSTEMS*, MACHINE LEARNING-POWERED ARTIFICIAL INTELLIGENCE IN ARMS CONTROL, Peace Research Institute Frankfurt
- Robert Mazzolin (2020), *Artificial Intelligence and Keeping Humans "in the Loop"*, MODERN CONFLICT AND ARTIFICIAL INTELLIGENCE, Centre for International Governance Innovation
- Rona, Gabor. (2026). *The Iranian School Strike: Excusable? A Violation? A War Crime?* *Opinio Juris*, 30 March 2026.
- Rudy Guyonneau, Arnaud Le Dez (2019), *Artificial Intelligence in Digital Warfare: Introducing the Concept of the Cyberteammate*, The Cyber Defense Review, Vol. 4, No. 2, pp. 103-116

- Santoni de Sio, Filippo & Van den Hoven, Jeroen. (2018). "Meaningful Human Control over Autonomous Systems: A Philosophical Account." *Frontiers in Robotics and AI*, Vol. 5, Article 15.
- Shai Farber (2025), *AI-Enabled Terrorism: A Strategic Analysis of Emerging Threats and Countermeasures in Global Security*, *Journal of Strategic Security*, Vol. 18, No. 3, pp. 320-344
- Ulrike Esther Franke (2019), *NOT SMART ENOUGH: THE POVERTY OF EUROPEAN MILITARY THINKING ON ARTIFICIAL INTELLIGENCE*, European Council on Foreign Relations
- VINCENT BOULANIN, LORA SAALMAN, PETR TOPYCHKANOV, FEI SU, MOA PELDAN CARLSSON (2020), *AI and the military modernization plans of nuclear-armed states, Artificial Intelligence, Strategic Stability and Nuclear Risk*, Stockholm International Peace Research Institute
- VLADISLAV CHERNAVSKIKH, JULES PALAYER (2025), *IMPACT OF MILITARY ARTIFICIAL INTELLIGENCE ON NUCLEAR ESCALATION RISK*, Stockholm International Peace Research Institute
- YANITRA KUMARAGURU (2020), *A pre-emptive ban on lethal autonomous weapon systems*, *The Impact of Artificial Intelligence on Strategic Stability and Nuclear Risk: Volume III South Asian Perspectives*, Stockholm International Peace Research Institute

Reports

- International Committee of the Red Cross (ICRC). (2019). *Autonomous Weapon Systems: Implications of Increasing Autonomy in the Critical Functions of Weapons*. Geneva: ICRC.
- International Committee of the Red Cross (ICRC). (2021). *Position on Autonomous Weapon Systems*. Geneva: ICRC.