

مسئولیت کیفری فرماندهان در عصر جنگ الگوریتمی؛ بازخوانی معیارهای «علم»، «کنترل مؤثر» و «پیشگیری» در پرتو حقوق بین‌الملل کیفری

پوریا احمدی

دانش‌آموخته کارشناسی ارشد حقوق جزا و جرم‌شناسی، واحد سمنان، دانشگاه آزاد اسلامی، سمنان، ایران.

poriaahmadi23@gmail.com

امین علی‌زاده (نویسنده مسئول)

دانشجوی دکتری حقوق کیفری و جرم‌شناسی، دانشکده حقوق و علوم سیاسی، پردیس دانشگاه تهران، تهران، ایران.

amin.alizadeh@srbiau.ac.ir

چکیده

جنگ الگوریتمی، دگرگونی بنیادینی در جایگاه تصمیم‌گیری در میدان نبرد پدید آورده است؛ تصمیمی که روزگاری در ذهن و اراده یک فرمانده انسانی شکل می‌گرفت، امروز به طور فزاینده‌ای در درون سامانه‌های خودکار و از طریق فرایندهای محاسباتی غیرقابل پیش‌بینی تولید می‌شود. این جابه‌جایی «کانون تصمیم» پرسشی حقوقی را بر می‌تابد که در دل آن بی‌توجهی به آن، معنای مسئولیت کیفری را در عرصه بین‌المللی به خطر می‌اندازد؛ در نظامی که فرمانده دیگر فرمان صریح صادر نمی‌کند، بلکه چارچوب‌هایی را تعریف می‌کند که درون آنها یک سامانه، مستقلانه تصمیم می‌گیرد، تکلیف انتساب مسئولیت به او چیست؟ مقاله حاضر با اتخاذ روش توصیفی-تحلیلی در پی پاسخ به این پرسش است که آیا چارچوب مسئولیت کیفری فرماندهان در ماده ۲۸ اساسنامه رم می‌تواند در برابر چالش‌های جنگ الگوریتمی، ظرفیت تفسیری خود را حفظ کند یا در برابر واقعیت سامانه‌های خودمختار فرو می‌پاشد. یافته اصلی پژوهش آن است که «گسست کنترلی» ناشی از خودمختاری سامانه‌ها، ماده ۲۸ را نه منسوخ، که به دیالکتیک جدی با شرایط نوین فرامی‌خواند؛ اما این دیالکتیک مستلزم بازخوانی انتقادی سه معیار بنیادین مسئولیت است؛ گذار از «علم قطعی» به «علم مبتنی بر ریسک»، صورت‌بندی مجدد «کنترل مؤثر» در قامت «کنترل سیستمیک» و بازتعریف «پیشگیری» به مثابه تکلیف نظارتی مستمر و مداخله انسانی معنادار. مقاله در نهایت مدل تحلیلی پیشنهاد می‌کند که با حفظ متن ماده ۲۸ و بدون توسعه بی‌ضابطه مسئولیت کیفری، عاملیت انسانی را در قلب حقوق بین‌الملل کیفری زنده نگه می‌دارد و در عین حال به قضات و دادستان‌های بین‌المللی راهنمایی عملی برای احراز مسئولیت در میدان‌های جنگ الگوریتمی ارائه می‌دهد. واژگان کلیدی: جنگ الگوریتمی، مسئولیت کیفری فرماندهان، ماده ۲۸ اساسنامه رم، کنترل مؤثر، کنترل انسانی معنادار، عاملیت انسانی.

جنگ، از دیرباز به مثابه میدانی برای آفرینش و به کارگیری فناوری‌های نوپدید شناخته شده است؛ اما هیچ‌یک از تحولات فنی پیشین، به اندازه پدیدار «جنگ الگوریتمی» به بنیاد حقوق بین‌الملل کیفری نزدیک نشده است. تحولات پیشین در ادوات جنگی از باروت تا جنگ‌افزارهای هسته‌ای ذیل ثابتی را بر هم نمی‌زدند؛ فرمانده انسانی همچنان تصمیم‌گیرنده‌ای بود که در لحظه حساس، اراده خویش را بر روند عملیات تحمیل می‌کرد. آنچه جنگ الگوریتمی به چالش می‌کشد، نه صرفاً ابزار نبرد، که جایگاه خود تصمیم است؛ جایی که روزگاری «اینجا و از سوی او» گرفته می‌شد، اکنون در درون سامانه‌های خودکار، از طریق موتورهای محاسباتی یادگیری‌محور و در مقیاسی غیرقابل ردیابی برای ذهن انسانی تولید می‌شود. این جابه‌جایی، با اصطلاحی که در ادبیات نوین حقوق بشردوستانه رواج یافته، «جابجایی کانون تصمیم»^۱ نامیده می‌شود (Emily & Pert, 2024: 261). این واقعیت نوپدید، نظام مسئولیت کیفری را در نقطه‌ای حساس مورد پرسش قرار می‌دهد. حقوق کیفری، پیش‌فرض مسلم خود را بر وجود عامل انسانی اراده‌مندی بنیان نهاده است که بتواند نتیجه رفتار خویش را پیش‌بینی کرده و از آن پیشگیری کند؛ و مسئولیت موسوم به «مسئولیت کیفری فرماندهان» در حقوق بین‌الملل کیفری، شاید برجسته‌ترین تجلی همین پیش‌فرض باشد. در واقع، ماده ۲۸ اساسنامه رم دیوان کیفری بین‌المللی، چنان که در همین نوشتار به تفصیل خواهیم دید، بر استعلایی‌ترین نقطه اتصال میان «قدرت فرماندهی» و «تکلیف انسانی» استوار است؛ فرمانده نه به‌خاطر آنچه خود انجام داده، بلکه به‌خاطر آنچه «نمی‌دانسته‌است» و یا «نگفته‌است» پاسخ‌گوست. دکترین مسئولیت فرماندهی، در حقیقت، برای گشودن سیاهچاله‌ای به کار گرفته شده است که در آن، سلسله‌مراتب فرماندهی، امکان پنهان‌شدن فرد پاسخ‌گو را فراهم می‌آورد؛ اما پرسش آن است که وقتی نه ابزار انجام، بلکه خود «تصمیم» نیز از دستان فرمانده انسانی خارج می‌شود، آیا این دکترین همچنان می‌تواند رسالت خود را ایفا کند؟

شکی نیست که مسئله مذکور دیگر در قلمرو آینده‌نگری صرف باقی نمانده و به واقعیتی ملموس در منازعات معاصر تبدیل شده است. سامانه‌های تسلیحاتی مستقل^۲ و ابزارهای تصمیم‌گیری الگوریتمی، نه در آزمایشگاه‌ها که در میدان‌های جنگ واقعی، در کنش‌های هدف‌گیری، شناسایی و حتی انتخاب گزینه‌های عملیاتی به کار گرفته شده‌اند. با وجود این، حقوق بین‌الملل با یک معاهده جامع برای سامان‌دهی این فناوری مواجه نیست؛ به‌طور نهایی در مورد جنگ افزارهای خودمختار، توافق در چارچوب کنوانسیون سلاح‌های متعارف، همچنان در سطح گفتگو و کارشناسی باقی مانده است. در این خلأ تقنینی، بر دوش حقوق کیفری بین‌الملل است که به‌عنوان «شبکه نهایی ایمنی»، از فروپاشی مفهوم پاسخگویی انسانی جلوگیری کند. این نوشتار در همین نقطه، منزلت حقوقی خود را تعریف می‌کند؛ تلاش برای پاسخ به این پرسش که آیا ماده ۲۸ اساسنامه رم، در برابر واقعیت سامانه‌های خودمختار، «ظرفیت تفسیری» خود را حفظ می‌کند یا در برابر این واقعیت فرو می‌پاشد؟

فرضیه این پژوهش، راه سوم و میانه‌روی میان دو پاسخ دوقطبی موجود است. در یک سو، رویکردی که با توسل به «خودمختاری» سامانه‌ها، اعمال مسئولیت بر فرمانده را ناممکن یا بی‌معنا می‌شمارد و به نوعی «مسئولیت بدون عامل» ختم می‌شود؛ و در سوی دیگر، رویکردی که با تفسیر منعطف و توسعه‌گرایانه از ماده ۲۸، فرمانده را در برابر هر خطای سامانه پاسخ‌گو می‌داند و بدین ترتیب، خطر «مسئولیت کیفری بی‌ضابطه» را در پی دارد. فرضیه نوشتار حاضر آن است که ماده ۲۸ اساسنامه رم، نه منسوخ و نه دچار فروپاشی شدن، بلکه «در دیالکتیک با شرایط نوین» قرار می‌گیرد؛ به این معنا که حفظ کلیدواژه‌های این ماده (علم، کنترل مؤثر و پیشگیری) با «بازخوانی انتقادی» آن‌ها در پرتو واقعیت فناورانه امکان‌پذیر است. این بازخوانی سه وجهی گذار از «علم قطعی» به «علم مبتنی بر ریسک»، صورت‌بندی مجدد «کنترل مؤثر» در قامت «کنترل سیستمیک»، و بازتعریف «پیشگیری» به مثابه «تکلیف نظارتی مستمر» به همراه «مداخله انسانی معنادار»، نه تنها دکترین مسئولیت فرماندهی را از فروپاشی نجات می‌دهد، بلکه به‌طور هم‌زمان، از توسعه بی‌ضابطه آن نیز جلوگیری می‌کند.

^۱ - displacement of the locus of decision

^۲ - Autonomous Weapon Systems

پژوهش پیش رو با اتخاذ روش توصیفی-تحلیلی، نخست به مفهوم‌شناسی جنگ الگوریتمی و تمایز آن از جنگ‌های کلاسیک پرداخته و در ادامه، حدود مفهوم «مسئولیت کیفری فرماندهان» و جایگاه آن در ماده ۲۸ اساسنامه رم را تبیین می‌کند. سپس گسست کنترلی ناشی از خودمختاری سامانه‌ها در ترازوی این دکترین گذارده می‌شود تا عمق چالش آشکار گردد. بخش بعدی که «بطن تحلیلی» نوشتار است، به بازخوانی انتقادی سه معیار بنیادین مسئولیت (علم، کنترل مؤثر و پیشگیری) اختصاص می‌یابد و در نهایت، بخش آخر، با ارائه مدل «مسئولیت‌پذیری انسانی»^۱ و چارچوب «کنترل انسانی معنادار»، راهکاری برای عبور از بحران انتساب ارائه می‌دهد.

۱- مفاهیم نظری

۱-۱- مفهوم جنگ الگوریتمی و تفاوت آن با جنگ‌های کلاسیک

جنگ الگوریتمی را می‌توان مرحله‌ای از دگرگونی جنگ دانست که در آن، بخش مهمی از شناسایی، اولویت‌بندی، هدف‌گیری، تخصیص منابع و حتی پیشنهاد یا اجرای واکنش عملیاتی، از قضاوت مستقیم انسانی به سامانه‌های مبتنی بر داده، مدل‌های یادگیری ماشین و فرایندهای خودکار واگذار می‌شود (Kania, 2017: 5-9). در این معنا، «الگوریتمی» بودن جنگ صرفاً به استفاده از فناوری دیجیتال یا پهپاد محدود نیست، بلکه ناظر بر انتقال تدریجی مرکز ثقل تصمیم‌گیری از انسان به معماری محاسباتی است؛ معماری‌ای که می‌تواند بر پایه داده‌های انبوه، الگوهای احتمالاتی و پردازش بلادرنگ، پیشنهادهایی بدهد یا کنش‌هایی را بدون مداخله فوری انسان انجام دهد.

این تحول از آن جهت اهمیت حقوقی دارد که جنگ کلاسیک، حتی در پیچیده‌ترین صورت‌های خود، بر یک زنجیره فرماندهی انسانی متکی بود. در جنگ کلاسیک، تصمیم برای حمله، انتخاب هدف، زمان‌بندی، تناسب، احتیاط و توقف عملیات، نهایتاً به اراده و ادراک یک مقام انسانی قابل انتساب بود؛ به همین دلیل، حقوق بین‌الملل بشردوستانه و حقوق بین‌الملل کیفری توانستند مفاهیمی مانند تفکیک، تناسب، احتیاطات لازم و مسئولیت فرماندهی را بر مبنای عاملیت انسانی صورت‌بندی کنند (Schmitt & Thurnher, 2013: 233-240). اما در جنگ الگوریتمی، این زنجیره به‌طور کامل از بین نمی‌رود، بلکه دچار جابه‌جایی می‌شود؛ انسان همچنان طراح، مجوزدهنده یا ناظر است، اما تصمیم عملیاتی در سطحی تولید می‌شود که برای او همیشه قابل پیش‌بینی، قابل توضیح یا حتی قابل مهار نیست (Payne, 2021: 33-39).

تفاوت بنیادی جنگ الگوریتمی با جنگ‌های کلاسیک در همین نقطه است. در جنگ کلاسیک، خطای انسانی را می‌توان با قواعد متعارف مسئولیت، از جمله علم، تقصیر، کنترل و امتناع از پیشگیری، تحلیل کرد؛ اما در جنگ الگوریتمی، مسئله فقط خطا نیست، بلکه «فاصله میان طراحی و نتیجه» است. سامانه‌های خودمختار یا نیمه‌خودمختار ممکن است در چارچوبی عمل کنند که توسعه‌دهنده یا فرمانده، همه پیامدهای آن را در لحظه آغاز عملیات نداند یا نتواند به‌طور دقیق کنترل کند (ICRC, 2018: 5-8). از این‌رو، جنگ الگوریتمی منطق کلاسیک جنگ را که بر مشاهده‌پذیری تصمیم انسانی و قابلیت انتساب مستقیم استوار بود، به چالش می‌کشد و یک لایه میانی از «پردازش ماشینی تصمیم» ایجاد می‌کند (Scharre, 2018: 57-63). از منظر حقوقی، این تفاوت پیامد مهمی دارد. در جنگ کلاسیک، مسئولیت عمدتاً بر پایه «اعمال یا ترک فعل انسان» سنجیده می‌شد؛ در جنگ الگوریتمی، پرسش اصلی این است که آیا انسان هنوز آن‌قدر در حلقه تصمیم باقی مانده که بتوان او را مسئول دانست یا نه. به بیان دقیق‌تر، اگر در جنگ سنتی «فرمان» منشأ مستقیم اثر عملیاتی بود، در جنگ الگوریتمی «چارچوب تصمیم» جای فرمان را می‌گیرد؛ و همین انتقال، مرزهای مسئولیت کیفری فرماندهان را مبهم‌تر می‌کند (Liu, 2012: 640-646). بنابراین، جنگ الگوریتمی نه صرفاً جنگی با ابزارهای نو، بلکه جنگی با منطق تازه در تولید تصمیم است. اگر جنگ کلاسیک بر «انسان تصمیم‌گیرنده» بنا شده بود، جنگ الگوریتمی بر «انسان طراح، ناظر یا مجازدهنده» استوار است؛ و همین تفاوت، در ادامه مقاله، مبنای بازخوانی مسئولیت کیفری فرماندهان در پرتو ماده ۲۸ اساسنامه رم خواهد بود.

¹ - human accountability

برای آنکه مفهوم جنگ الگوریتمی از یک انتزاع نظری خارج شود و ابعاد عملیاتی و حقوقی آن عینیت یابد، سه نمونه مستند از منازعات اخیر را می‌توان بررسی کرد؛ به‌کارگیری سامانه‌های خودمختار در جنگ لیبی (۲۰۲۰-۲۰۲۱)، استفاده گسترده از سامانه‌های تصمیم‌یار هوشمند در جنگ غزه (۲۰۲۳-۲۰۲۴)، و افزایش وابستگی به مهمات پرسه‌زن و پردازش خودکار هدف در جنگ اوکراین.

الف) سامانه‌های خودمختار در لیبی و نخستین گزارش رسمی سازمان ملل: مستندترین نمونه اولیه جنگ الگوریتمی در میدان نبرد، استفاده از پهپاد انتحاری خودمختار ترکیه‌ای «کارگو-۲»^۱ در جریان درگیری‌های لیبی است. گزارش هیئت کارشناسان سازمان ملل متحد در مورد لیبی (سند مورخ ۸ مارس ۲۰۲۱)^۲ به‌صراحت گزارش داده است که این سامانه قابلیت «شلیک، فراموش کردن و یافتن»^۳ هدف را دارد؛ یعنی پهپاد پس از شناسایی هدف، بدون نیاز به فرمان مجدد انسانی، آن را رهگیری و مورد حمله قرار می‌دهد (United Nations Security Council, 2021: 17). این گزارش، نخستین تأیید رسمی در سطح سازمان ملل از به‌کارگیری واقعی یک سامانه تسلیحاتی دارای عملکرد خودمختار در کشتار بود و نشان داد که بحث از حوزه نظری به تجربه میدانی رسیده است (Bátori, 2025: 2-4). اهمیت حقوقی این نمونه در آن است که «انتخاب و رهگیری هدف» به‌طور کامل از حلقه کنترل انسانی خارج شده بود؛ امری که پیش‌تر تنها در مباحث هنجاری مطرح می‌شد.

ب) سامانه‌های تصمیم‌یار هوشمند در جنگ غزه: گزارش‌های مشترک^۴ و بازتاب‌های آن در رسانه‌های معتبر، از دو سامانه هوش مصنوعی با نام‌های «لاوندر»^۵ و «گاسپل»^۶ پرده برداشتند که در عملیات اسرائیل در غزه به‌کار گرفته شدند. بنا بر این گزارش‌ها، لاوندر فهرستی از ده‌ها هزار فرد به‌عنوان «عضو مشکوک حماس» تولید می‌کرد و گاسپل نیز پیشنهادهای خودکار برای حمله به اماکن مسکونی و زیرساخت‌ها ارائه می‌داد (Abraham, 2024). آنچه این نمونه را به‌ویژه برای بحث مسئولیت فرماندهی مهم می‌کند، ادعای خود منابع و تحلیلگران حقوقی است؛ به گفته شش افسر اطلاعاتی، تأثیر خروجی سامانه بر عملیات چنان بود که آن را «همچون تصمیم انسانی» تلقی می‌کردند و در برخی موارد، تأیید انسانی صرفاً به یک «مهر لاستیکی» تبدیل شده بود که گاهی در حدود ۲۰ ثانیه انجام می‌شد (Goodwin, 2024). در همین راستا، بررسی‌ها بر این نکته تأکید کرده‌اند که اتکای گسترده به این سامانه‌ها، دشواری‌های جدی در ایفای تعهدات احتیاطی^۷ و تفکیک^۸ ایجاد کرده و سرعت و وسعت هدف‌گیری را چنان افزایش داده است که کنترل‌پذیری انسانی را به‌شدت کاهش می‌دهد (Andersin, 2025: 336-352) (Blank, 2024). این نمونه، برخلاف کارگو-۲، نه بر خودمختاری در اجرای حمله، بلکه بر خودمختاری تصمیم‌یاری در انتخاب و اولویت‌بندی اهداف استوار است و نشان می‌دهد جنگ الگوریتمی صرفاً به «سلاح خودمختار» محدود نیست؛ بلکه به سامانه‌های پشتیبانی تصمیم نیز گسترش یافته است که حلقه انسانی را به جایگاه ناظر تأییدکننده تنزل می‌دهند.

ج) پردازش خودکار هدف و مهمات پرسه‌زن در جنگ اوکراین: جنگ اوکراین نیز شواهدی از گذار به جنگ الگوریتمی ارائه می‌دهد، هرچند با درجه‌ای متفاوت از دو نمونه پیشین. در این منازعه، استفاده از مهمات پرسه‌زن^۹ و سامانه‌های شناسایی و هدف‌گیری مبتنی بر یادگیری ماشین به‌صورت گسترده گزارش شده است؛ سامانه‌هایی که توانایی رهگیری خودکار هدف را دارند و در بخش‌هایی از زنجیره «شناسایی، انتخاب و رهگیری»، وابستگی به کنترل لحظه‌به‌لحظه انسانی را کاهش می‌دهند (SIPRI, 2024) (Horowitz & Scharre, 2023). همچنین گزارش‌های تحلیلی نشان می‌دهند که نرخ هدف‌گیری مؤثر سامانه‌های مجهز به

¹ - Kargu-2

² - S/2021/229

³ - fire, forget and find

⁴ - +972 Magazine/Local Call

⁵ - Lavender

⁶ - Gospel

⁷ - precaution

⁸ - distinction

⁹ - loitering munitions

هوش مصنوعی در مقایسه با ابزارهای دستی سنتی به طور چشمگیری افزایش یافته است (Breaking Defense, 2024). اگرچه اوکراین نمونه‌ای «تمام خودمختار» نیست، اما تجربه آن نشان می‌دهد که جنگ الگوریتمی بیش از آنکه یک وضعیت قطعی باشد، یک طیف است از خودمختاری کامل در انتخاب و حمله (کارگو-۲) تا خودمختاری در تصمیم‌یاری و فیلتر اهداف (لاوندر و گاسپل) و در نهایت خودمختاری جزئی در رهگیری (نمونه اوکراین). در همه این مراتب، نقطه مشترک یکسان است؛ یعنی جابه‌جایی کانون تصمیم از انسان به معماری محاسباتی، که همان معیار تمایز جنگ الگوریتمی از جنگ کلاسیک است. این سه نمونه در مجموع نشان می‌دهند که جنگ الگوریتمی دیگر یک فرضیه درباره آینده نیست، بلکه یک واقعیت جاری در میدان‌های نبرد معاصر است؛ و همین واقعیت، پرسش از معیارهای سنتی مسئولیت کیفری فرماندهان را از یک بحث انتزاعی به یک ضرورت عملی بدل می‌کند.

۱-۲- مفهوم و حدود «مسئولیت کیفری فرماندهان» در حقوق بین‌الملل کیفری

دکترین مسئولیت فرماندهی یا مافوق، به عنوان یکی از قواعد بنیادین حقوق بین‌الملل کیفری عرفی، سازوکاری حقوقی برای انتساب مسئولیت کیفری به فرماندهان و مافوق‌های نظامی و غیرنظامی به دلیل ترک فعل^۱ در قبال جرایم ارتكابی توسط نیروهای تحت امر آنان است (Cassese, 2013: 215). این دکترین که ریشه‌های تاریخی آن به اصول حقوق جنگ و دادگاه‌های نظامی پس از جنگ جهانی دوم بازمی‌گردد، در جریان محاکمه ژنرال یاماشیتا توسط کمیسیون نظامی ایالات متحده متجلی شد؛ جایی که دادگاه اعلام کرد وظیفه فرماندهی مستلزم اتخاذ تدابیر مقتضی جهت کنترل نیروها برای جلوگیری از نقض قوانین جنگ است (Yamashita Case, 1946: 43-44). تکامل قضایی این مفهوم در جریان دادگاه‌های نظامی نورنبرگ و دادگاه نظامی بین‌المللی شرق دور تداوم یافت و در نهایت، پرونده هادامار^۲ و پرونده فرماندهی عالی^۳ چارچوب‌های اولیه را برای سنجش «وظیفه مراقبت و اعمال نظارت» فرمانده تثبیت کردند (Ambos, 2014: 181-183). با این حال، تبلور رسمی و نظام‌مند این قاعده به عنوان یک تعهد حقوقی مستقل، ابتدا در بند ۲ ماده ۸۶ و ماده ۸۷ پروتکل الحاقی اول به کنوانسیون‌های ژنو ۱۹۴۹ صورت پذیرفت و متعاقباً در اساسنامه دادگاه‌های کیفری اختصاصی برای یوگسلاوی سابق^۴ و رواندا^۵ گنجانده شد. اساسنامه رم در بندهای (الف) و (ب) ماده ۲۸ خود، با تفکیک میان مافوق‌های نظامی و غیرنظامی، این دکترین را به اوج تکامل هنجاری خود رساند (Schabas, 2016: 448). تحلیل ساختاری مسئولیت فرماندهی در رویه قضایی بین‌المللی نشان می‌دهد که انتساب این مسئولیت نیازمند احراز سه عنصر مادی و روانی به هم پیوسته است که نادیده گرفتن هر یک از آن‌ها، زنجیره علیت حقوقی را قطع می‌کند. عنصر نخست، وجود رابطه مافوق-مادون بر پایه معیار «کنترل مؤثر»^۶ است. شعبه محاکمه دادگاه یوگسلاوی سابق در پرونده تاریخی سلیبیچی^۷ تصریح کرد که کنترل مؤثر به معنای توانایی مادی فرمانده برای «پیشگیری از ارتکاب جرایم» یا «مجازات مرتکبان» است، فارغ از اینکه این اقتدار به شکل قانونی^۸ باشد یا عملاً^۹ اعمال شود (ICTY, 1998: para 378). دیوان بین‌المللی کیفری نیز در اولین پرونده مرتبط با ماده ۲۸ یعنی پرونده بمبا،^{۱۰} ضمن تأیید این معیار، استدلال کرد که کنترل مؤثر باید فراتر از توانایی نظری برای صدور دستور باشد و به توانایی واقعی فرمانده در هدایت رفتار نیروها و مهار رفتارهای مجرمانه آن‌ها اشاره داشته باشد (ICC, 2016: para 168). عنصر دوم، رکن روانی یا آگاهی فرمانده است که در اساسنامه رم با فرمول‌بندی «می‌داند یا بایستی می‌دانست»^{۱۱} تعریف شده است. رویه قضایی سنتی در دادگاه یوگسلاوی سابق، اعمال معیار «باید می‌دانست» را منوط به وجود اطلاعاتی می‌کرد که فرمانده

¹- Omission

²- Hadamar Case

³- High Command Case

⁴- ICTY

⁵- ICTR

⁶- Effective Control

⁷- Čelebići Case

⁸- De Jure

⁹- De Facto

¹⁰- Bemba Case

¹¹- Know or should have known

را در معرض ظن قرار دهد (Hadžihasanović Case, 2008: para 28). با این حال، ماده ۲۸ اساسنامه رم استاندارد سخت‌گیرانه‌تر را برای فرماندهان نظامی وضع کرده است؛ به این معنا که فرمانده به دلیل وظیفه فعال نظارتی خود، نمی‌تواند به جهل داوطلبانه استناد کند و موظف است سیستم‌های کسب اطلاعات و گزارش‌دهی دقیقی را به کار گیرد تا از بروز جنایات مطلع شود (Ambos, 2014: 192).

عنصر سوم، عنصر مادی ترک فعل است که عبارت است از قصور فرمانده در اتخاذ «تمام تدابیر لازم و معقول» که در حیطه اختیارات مادی او برای جلوگیری از وقوع جرم، متوقف ساختن آن یا ارجاع موضوع به مقامات صالح برای تعقیب قضایی قرار دارد. شعبه تجدیدنظر دیوان بین‌المللی کیفری در پرونده بمبا شفاف‌سازی کرد که «لازم و معقول بودن» اقدامات فرمانده نباید بر اساس یک استاندارد انتزاعی ارزیابی شود، بلکه باید با توجه به توانایی مادی واقعی فرمانده در شرایط خاص زمانی و مکانی سنجیده شود (ICC, 2018: para 170). از این‌رو، مسئولیت فرماندهی در حقوق بین‌الملل کیفری نه بر مسئولیت عینی، بلکه بر مسئولیت ناشی از تقصیر در انجام وظیفه نظارتی^۱ استوار است.

۲- گسست کنترل: تلاقی دکترین مسئولیت فرماندهی با سیستم‌های خودمختار در پرتو ماده ۲۸ اساسنامه رم

هسته‌ی دکترین مسئولیت فرماندهی، در هر دو شکل نظامی و غیرنظامی آن که در ماده ۲۸ اساسنامه رم تبلور یافته، بر یک پیش‌فرض معرفت‌شناختی و عملی بنا شده است. فرمانده، به‌عنوان نقطه‌ی کانونی اراده و نظارت، در فاصله‌ای معنادار از صحنه‌ی ارتکاب جرم قرار دارد و از همین فاصله است که می‌تواند بر رفتار نیروهای تحت امر خود «اثر بگذارد». ماده ۲۸ این پیش‌فرض را در قالب سه معیار به‌هم‌پیوسته‌ی صورت‌بندی کرده است؛ وجود رابطه‌ی مبتنی بر «کنترل مؤثر» میان مافوق و مادون (بند الف)، احراز رکن روانی «علم» به‌معنای «می‌داند یا بایستی می‌دانست» (بند الف-۱)، و سرانجام تکلیف به اتخاذ «تمام تدابیر لازم و معقول» برای پیشگیری و سرکوب (بند الف-۲). تمامی این معیارها، به‌صورت ضمنی اما قاطع، بر یک الگوی رویدادمحور و انسانی استوارند؛ جرم، توسط کنشگرانی انسانی و قابل مشاهده ارتکاب می‌یابد و فرمانده، از طریق زنجیره‌ای از فرمان، نظارت و انضباط که همچنان در اختیار اراده‌ی اوست، می‌تواند در این زنجیره مداخله کند. اما وقتی فناوری خودمختار وارد میدان می‌شود، همین پیش‌فرض‌های اساسی فرو می‌ریزند و گسستی عمیق در قلب ماده ۲۸ ایجاد می‌شود که در ادامه، هر یک از این معیارها را به‌طور منفرد و در بستر همین گسست مورد واکاوی قرار می‌دهیم.

بحران در بستر «کنترل مؤثر»؛ نخستین و بنیادی‌ترین گسست، به معیار «کنترل مؤثر» مربوط می‌شود. شعبه محاکمه دادگاه یوگسلاوی سابق در پرونده سلبیچی این معیار را در قالب توانایی فرمانده برای «پیشگیری از ارتکاب» یا «مجازات مرتکبان» تعریف کرد (ICTY, 1998: para 378)، و شعبه تجدیدنظر دیوان بین‌المللی کیفری در پرونده بمبا نیز بر بُعد «مادی» این توانایی، یعنی وجود اقتدار واقعی و نه صرفاً صوری برای مهار رفتار نیروها، تأکید نهاد (ICC, 2018: para 170). این معیار در قلمرو نیروهای انسانی معنادار است، زیرا فرمانده در لحظات حساس می‌تواند با صدور فرمان، عزل، یا مداخله‌ی فیزیکی، زنجیره‌ی رفتاری سرباز را قطع کند. اما در جنگ الگوریتمی، منطق عملیاتی بر پایه‌ی «سرعت ماشین»^۲ بنا شده است؛ سامانه‌هایی که در کسری از ثانیه، بر پایه‌ی پردازش انبوه داده‌ها، هدف‌گیری و شلیک را انجام می‌دهند، دیگر در حصار زمان انسانی فرمانده قرار نمی‌گیرند تا او بتواند در آنها مداخله‌ی بلادرنگ کند. گزارش‌های میدانی از کاربرد پهپادهای خودمختار شناسایی شده در درگیری‌های لیبی و سامانه‌های هدف‌گیر هوش مصنوعی در غزه نشان داده‌اند که سرعت و استقلال تصمیم این سامانه‌ها به‌گونه‌ای است که امکان «بازبینی انسانی» در لحظه‌ی هدف‌گیری را عملاً منتفی می‌سازد (UNSC, 2021: 17). در چنین شرایطی، این پرسش حیاتی مطرح می‌شود که آیا فرمانده‌ای که فرمان استقرار چنین سامانه‌ای را صادر کرده، در لحظه‌ی ارتکاب جنایت ناشی از خطای الگوریتم، همچنان دارای «کنترل مؤثر» بر رفتار عامل جرم محسوب می‌شود یا اینکه این گسست زمانی و شناختی، رابطه‌ی مافوق-مادون را به حدی فرسوده می‌کند که انتساب مسئولیت بر مبنای ماده ۲۸ را ناممکن می‌سازد.

¹ - Fault-based liability for omission

² - machine-speed

بحران در بستر «علم»؛ گسست دوم، به رکن روانی ماده ۲۸ و معیار «می‌داند یا بایستی می‌دانست» بازمی‌گردد. رویه‌ی قضایی بین‌المللی این معیار را بر پایه‌ی وجود اطلاعاتی می‌سنجد که فرمانده را در وضعیت ظن نسبت به وقوع جرم قرار می‌دهد (ICTY, 2008: para 28). این منطق بر این فرض استوار است که رفتار عامل جرم، در قلمرو شناخت انسانی فرمانده قرار می‌گیرد و او قادر است نشانه‌های هشداردهنده را درک و ارزیابی کند. اما رفتار سامانه‌ی خودمختار ذاتاً با «ابهام محاسباتی» و پدیده‌ی «جعبه سیاه»^۱ آمیخته است؛ به این معنا که فرآیند درونی تصمیم‌گیری الگوریتم، حتی برای طراحان آن نیز به‌طور کامل شفاف و قابل تبیین نیست و پیش‌بینی دقیق خطاهای آن خارج از حیطه‌ی توان شناختی انسان قرار می‌گیرد (Chengeta, 2016: 39). در چنین فضایی، پرسش این است که آیا جهل فرمانده نسبت به خطای پیش‌روی الگوریتم، جهلی «قابل توجیه» است که او را از بار «بایستی می‌دانست» می‌رهاند، یا اینکه «وظیفه‌ی مراقبت» در قلمرو فناوری‌های خودمختار، به‌گونه‌ای بازتعریف می‌شود که فقدان دانش فنی فرمانده خود به مثابه‌ی «قصور مستقل» و مبنایی برای انتساب مسئولیت عمل کند.

بحران در بستر «پیشگیری»؛ گسست سوم، در نهایت به تکلیف «اتخاذ تمام تدابیر لازم و معقول» مربوط می‌شود که در پرونده بمبا به‌عنوان الزامی بر پایه‌ی توانایی مادی فرمانده تفسیر شده است (ICC, 2016: para 168). در قلمرو نیروهای انسانی، این تکلیف با سازوکارهایی چون صدور فرمان، اعمال انضباط، پایش میدانی و ارجاع به مراجع قضایی محقق می‌شود؛ ابزارهایی که همگی در قلمرو اراده و اختیار مستمر فرمانده قرار دارند. اما در جنگ الگوریتمی، فرمانده در جایگاه «معمار عملیاتی» قرار می‌گیرد که توانایی‌اش دیگر نه به‌صورت واکنش‌های لحظه‌ای، بلکه به‌صورت طراحی پارامترها، محدودیت‌ها و قواعد تعامل با سامانه در مرحله‌ی پیش از استقرار متجلی می‌شود (Gunawan, et al, 2022: 9). پرسش بنیادین این است که آیا این نوع نوین «کنترل طراحی» را می‌توان با همان دامنه و آثار حقوقی «کنترل عملیاتی» در ماده ۲۸ هم‌ارز دانست و آیا تکلیف پیشگیری، فرمانده را موظف می‌کند تا سامانه‌هایی را که از «نظارت‌پذیری» کافی برخوردار نیستند، از استقرار باز دارد (ICRC, 2019: 12). به بیان دیگر، آیا در جهان خودمختار، «نظارت انسانی معنادار» به شرط امکان‌پذیری خود مسئولیت بدل می‌شود و غیاب آن، گسستی ایجاد می‌کند که ماده ۲۸ را در مقابل چالش «خلأ مسئولیت» قرار می‌دهد.

در مجموع، آنچه در این بخش بر آن انگشت نهاده شد، صرفاً یک دشواری فنی یا عملیاتی نیست، بلکه بحرانی هنجاری در قلب ماده ۲۸ است. هر سه معیار «کنترل مؤثر»، «علم» و «پیشگیری» بر منطق عاملیت انسانی مستمر بنا شده‌اند، در حالی که جنگ الگوریتمی همین منطق را با جابجایی کانون تصمیم از انسان به ماشین، از میان برداشته است. این گسست کنترل، خلأ مسئولیت کیفری را رقم می‌زند؛ خلائی که یا باید با بازخوانی انتقادی همین معیارها پر شود، یا آن‌که مسئولیت فرمانده در جهان جدید، به کلی در غبار جعبه‌ی سیاه گم شود. بخش آتی، یعنی بازخوانی انتقادی معیارهای سه‌گانه، دقیقاً به همین چالش پاسخ خواهد داد.

۳- بازخوانی انتقادی معیارهای سه‌گانه: بحران قابلیت انتساب؟

۳-۱- از «علم قطعی» به «علم مبتنی بر ریسک»

ماده ۲۸ اساسنامه رم، مسئولیت کیفری فرماندهان و سایر مافوق‌ها را بر مبنای یک منطق آگاهانه از نظارت، پیشگیری و واکنش بنا می‌کند. در این ماده، معیار ذهنی انتساب، برای فرمانده نظامی «می‌داند یا بایستی می‌دانست» است و برای مافوق غیرنظامی «می‌دانست یا آگاهانه از دریافت اطلاعاتی که به‌وضوح دلالت بر ارتکاب جرم داشت، چشم پوشید» آمده است. بنابراین، حتی در صورت فقدان اثبات علم مستقیم به هر فعل مجرمانه، نویسندگان سند، مسئولیت را به نقطه‌ای پیش‌تر منتقل کرده است؛ به مرحله‌ای که فرمانده، به‌واسطه جایگاه، اطلاعات در دسترس، و تکالیف نظارتی‌اش، باید از خطر ارتکاب جرم آگاه می‌بود. این ساختار نشان می‌دهد که ماده ۲۸ از همان ابتدا بر «علم قطعی» به معنای آگاهی کامل و بالفعل از وقوع جرم بنا نشده است. برعکس، هسته هنجاری آن بر یک استاندارد مبتنی بر ریسک استوار است. فرمانده مکلف است نه فقط وقوع بالفعل جرم، بلکه خطر معقول و قابل پیش‌بینی ارتکاب جرم را نیز در نظر بگیرد و در برابر آن اقدام کند. در این چارچوب، «بایستی می‌دانست» را

¹ - black box

نباید صرفاً به معنای یک آگاهی ذهنی ضعیف فهم کرد؛ این عبارت، متضمن یک تکلیف فعال برای دریافت، ارزیابی و واکنش به اطلاعات هشداردهنده است. به بیان دیگر، معیار ماده ۲۸، فرمانده را از سطح «دانستنِ پسینی» به سطح «مدیریتِ پیشینی خطر» منتقل می‌کند. این برداشت با منطق رسیدگی‌های دیوان کیفری بین‌المللی نیز سازگار است. در پرونده بمبا، مسئله اصلی دقیقاً این بود که آیا انتساب مسئولیت فرماندهی باید بر پایه آگاهی واقعی از تمام جزئیات اعمال زیردستان سنجیده شود یا بر پایه عدم انجام تدابیر لازم در برابر اطلاعاتی که می‌توانست فرمانده را به خطر ارتکاب جرم متوجه کند. شعبه تجدیدنظر دیوان یادآور شد که ماده ۲۸ به دنبال یک معیار سخت‌گیرانه اما پیش‌گیرانه است؛ معیاری که فرمانده را از پناه بردن به ادعای «من شخصاً صحنه جرم را ندیده‌ام» باز می‌دارد، اگر اطلاعات موجود و ساختار فرماندهی او را موظف به واکنش می‌کرد.

در جنگ الگوریتمی، اهمیت این بازخوانی دوجندان می‌شود. هنگامی که تصمیم‌گیری عملیاتی به سامانه‌های خودکار یا نیمه‌خودکار سپرده می‌شود، علم فرمانده دیگر لزوماً از گزارش مستقیم میدانی نمی‌آید؛ بلکه از ویژگی‌های فنی سامانه، نتایج آزمایش‌ها، نرخ خطا، محدودیت‌های محیط عملیاتی و هشدارهای متخصصان فنی به دست می‌آید. در چنین وضعی، اگر فرمانده از ریسک‌های شناخته‌شده سامانه، مانند ناتوانی در تشخیص دقیق هدف یا افزایش خطا در محیط‌های شهری، آگاه باشد یا باید آگاه می‌بود، نمی‌تواند با استناد به خودمختاری ماشین از مسئولیت بگریزد. آنچه اهمیت می‌یابد، نه «علم به نتیجه نهایی»، بلکه علم به ریسک قابل پیش‌بینی ناشی از به‌کارگیری سامانه است. از این منظر، «علم مبتنی بر ریسک» به‌مثابه تفسیر کارکردی ماده ۲۸، یک جابه‌جایی ناموجه در متن اساسنامه ایجاد نمی‌کند؛ بلکه ظرفیت نهفته در خود ماده را آشکار می‌سازد. زیرا منطق مسئولیت فرماندهی از آغاز بر این پیش‌فرض استوار بوده که مافوق، به سبب قدرت سازمانی و دسترسی به اطلاعات، باید خطر را شناسایی و مهار کند. تفاوت عصر جنگ الگوریتمی در این است که خطر دیگر صرفاً از رفتار پیش‌بینی‌نشده زیردست انسانی ناشی نمی‌شود، بلکه از معماری فنی سامانه و نحوه استقرار آن در میدان نبرد سرچشمه می‌گیرد. در نتیجه، معیار «باید می‌دانست» باید به‌گونه‌ای فهم شود که شامل آگاهی یا قابلیت آگاهی از ریسک سیستمی نیز باشد.

این خوانش، هم از حیث سیاست جنایی بین‌المللی و هم از حیث تفسیر نص، موجه‌تر از اصرار بر مفهوم مضیق «علم قطعی» است. اگر حقوق کیفری بین‌المللی بخواهد در برابر فناوری‌های خودمختار کارآمد بماند، باید مسئولیت فرمانده را به نقطه‌ای پیوند دهد که تصمیم انسانی برای استقرار، فعال‌سازی و اعتماد به سامانه گرفته می‌شود. به همین دلیل، در عصر جنگ الگوریتمی، معیار واقعی ماده ۲۸ نه «آیا فرمانده بعداً یقین پیدا کرد؟»، بلکه این است که آیا او پیشاپیش باید می‌فهمید که سامانه، در آن زمینه عملیاتی، ریسک غیرقابل قبول ارتکاب نقض شدید را ایجاد می‌کند یا نه.

۲-۳- «کنترل مؤثر» در قامت «کنترل سیستمیک»

در حقوق کیفری بین‌المللی، «کنترل مؤثر» در اصل برای توصیف رابطه‌ای به‌کار می‌رود که در آن مافوق، توان مادی جلوگیری از وقوع جرم، سرکوب آن، یا ارجاع موضوع به مقامات صلاحیت‌دار را دارد. ماده ۲۸ اساسنامه رم این منطق را به‌صراحت در قالب تکلیف فرمانده به اتخاذ «تمام تدابیر لازم و معقول در حدود قدرت خویش» بیان می‌کند. بنابراین، کنترل مؤثر هرگز صرفاً یک عنوان تشریفاتی یا جایگاه سازمانی نیست؛ بلکه به قدرت واقعی مداخله در فرآیند ارتکاب جرم وابسته است. دیوان کیفری بین‌المللی در پرونده بمبا نیز همین فهم را تقویت کرد، زیرا مسئولیت فرمانده را به این پرسش گره زد که آیا متهم از نظر مادی در موقعیتی بود که بتواند از جرم جلوگیری کند یا در برابر آن واکنش مؤثر نشان دهد.

اما در جنگ الگوریتمی، این معیار اگر به‌صورت ایستا فهم شود، بخشی از واقعیت قدرت را از دست می‌دهد. چراکه قدرت فرمانده در چنین محیطی دیگر فقط بر بدن‌های انسانی تحت امر اعمال نمی‌شود، بلکه بر زنجیره‌ای از انتخاب‌های فنی، طراحی نرم‌افزار، قواعد درگیری، آموزش سامانه، انتخاب داده‌های ورودی، محدودیت‌های عملیاتی و مجوز به‌کارگیری نیز اعمال می‌شود. به بیان دقیق‌تر، آنچه باید مورد سنجش قرار گیرد نه فقط این است که فرمانده مستقیماً بر سرباز یا اپراتور کنترل داشته یا نه، بلکه این است که آیا او بر معماری تصمیم‌گیری جنگ کنترل داشته است یا خیر. در این معنا، کنترل مؤثر از سطح کنترل فردی به سطح کنترل سیستمیک ارتقا پیدا می‌کند. این جابه‌جایی مفهومی به این معنا نیست که ماده ۲۸ نیازمند اصلاح باشد؛ بلکه باید فهمید

که در شرایط جدید، «قدرتِ مافوق» دیگر در یک نفر یا در یک لحظه متجلی نمی‌شود، بلکه در توان هدایت و محدودسازی سامانه‌ای که تصمیم را تولید می‌کند ظهور می‌یابد. اگر فرمانده سامانه‌ای را مستقر می‌کند که به‌طور ساختاری در تشخیص هدف دچار خطاست، یا اگر با علم به محدودیت‌های آن همچنان استفاده‌اش را در محیطی متراکم و غیرقابل تفکیک مجاز می‌کند، او در واقع از طریق انتخاب‌های پیشینی خود بر نتیجه عملیاتی کنترل اعمال کرده است. در این سطح، کنترل، مستقیماً با «فرمانِ شلیک» برابر نیست؛ بلکه به فرمانِ طراحی، استقرار، نظارت، محدودسازی و تعلیق گسترش می‌یابد.

ادبیات جدید نیز همین گره را برجسته کرده است. برخی نویسندگان، در بحث درباره سامانه‌های خودمختار، تأکید کرده‌اند که مسئولیت فرماندهی همچنان بر یک رابطه انسانی و سلسله‌مراتبی استوار است، اما «کنترل مؤثر» در این فضا باید به‌گونه‌ای فهم شود که شامل کنترل واقعی بر استفاده از سامانه و شرایط به‌کارگیری آن باشد، نه صرفاً رابطه مستقیم با عامل انسانی مجری (Sehrawat, 2020: 332). این نکته مهم است، زیرا اگر کنترل را فقط به تعامل فیزیکی با زیردست انسانی محدود کنیم، بخش مهمی از قدرت تصمیم‌گیری در جنگ‌های امروز از قلم می‌افتد. در عمل، بسیاری از تصمیم‌های مرگ‌بار در لایه‌ای بالاتر از صحنه نبرد اتخاذ می‌شوند؛ در سطح مجوز عملیاتی، تعریف مأموریت، تنظیم پارامترهای هدف‌گیری، و پذیرش ریسک. همین‌جا است که کنترل سیستمیک واقعیت پیدا می‌کند. از این منظر، «کنترل مؤثر» در عصر جنگ الگوریتمی باید به‌عنوان توان مؤثر برای مهار ریسک تولیدشده توسط سامانه بازخوانی شود. این توان ممکن است از طریق انتخاب یا رد سامانه، تنظیم محدودیت‌های عملکردی، الزام به نظارت انسانی معنادار، توقف عملیات در صورت بروز خطا، یا حتی عدم استقرار سامانه در محیط نامناسب اعمال شود. پس پرسش اصلی دیگر این نیست که آیا فرمانده در لحظه وقوع، دکمه را شخصاً فشار داده است یا نه؛ پرسش اصلی این است که آیا او در سطحی که تصمیم واقعی گرفته می‌شود، قدرت معنادار برای جلوگیری از پیامد مجرمانه را داشته است یا خیر. اگر پاسخ مثبت باشد، عنصر کنترل مؤثر برقرار است، هرچند کنترل مزبور از جنس کنترل بر الگوریتم، معماری عملیاتی و چرخه تصمیم باشد، نه صرفاً کنترل بر یک زیردست انسانی.

این بازخوانی، هم با متن ماده ۲۸ سازگار است و هم با هدف آن. هدف قاعده، فرار از مسئولیت در ساختارهای پیچیده نیست؛ هدف، جلوگیری از این است که مافوق با اتکا به لایه‌های واسط یا واسطه‌های فنی، از پاسخ‌گویی بگریزد. بنابراین، در جنگ الگوریتمی، کنترل مؤثر باید به‌صورت کنترل بر سیستم تصمیم‌ساز فهم شود. هر جا که فرمانده بتواند از طریق نظارت، محدودسازی یا تعلیق، از تولید خطر مجرمانه جلوگیری کند، اما چنین نکند، مسئولیت او در چارچوب ماده ۲۸ قابل طرح است.

۳-۳- «پیشگیری» به مثابه «تکلیفِ نظارتی»

تکلیف به پیشگیری^۱ در دکترین مسئولیت مافوق، نه یک تعهد به نتیجه، بلکه یک تعهد به وسیله است که بر مبنای اتخاذ «تمام تدابیر لازم و معقول»^۲ سنجیده می‌شود. در رویه قضایی محاکم بین‌المللی، به‌ویژه در پرونده اوریچ^۳ در دادگاه کیفری بین‌المللی یوگسلاوی سابق، تأکید شده است که آنچه یک تدبیر را «لازم و معقول» می‌سازد، بسته به شرایط عینی هر پرونده و قدرت واقعی فرمانده برای مداخله متفاوت است (ICTY, Orić, 2006, para. 338). این بدان معناست که پیشگیری صرفاً یک عمل منفعلانه نیست؛ بلکه مستلزم یک نظام نظارتی فعال است که پیش از وقوع هرگونه نقض، فعال شود.

در چارچوب جنگ‌های الگوریتمی، این تکلیفِ پیشگیرانه بازتعریف می‌شود. وقتی یک سامانه با خودمختاری عملیاتی در میدان نبرد حضور دارد، «پیشگیری» دیگر نمی‌تواند صرفاً به معنای صدور دستور توقف به یک سرباز در آستانه ارتکاب جرم باشد. در اینجا، پیشگیری با مفهوم نظارت سیستمیک گره می‌خورد. این نظارت نه بر رفتار فردی، بلکه بر صحت عملکرد، ثبات الگوریتمی و انطباق مستمر سامانه با قواعد حقوق بشردوستانه اعمال می‌شود. به عبارت دیگر، در محیط الگوریتمی، تکلیف به پیشگیری به

¹ - Duty to Prevent

² - all necessary and reasonable measures

³ - Orić

معنای تکلیف به تضمین کنترل انسانی معنادار^۱ در تمام مراحل چرخه هدف‌گیری است. بنابراین، «تدابیر لازم و معقول» برای یک فرمانده در عصر الگوریتم شامل طیفی از اقدامات پیشینی است؛ از حصول اطمینان نسبت به انجام تست‌های ارزیابی قانونی (موضوع ماده ۳۶ پروتکل الحاقی اول ۱۹۷۷) گرفته تا تنظیم دقیق پارامترهای عملیاتی و پایش مداوم سامانه برای شناسایی هرگونه انحراف یا خطای سیستماتیک. برخی تحلیل‌گران در بررسی دامنه مسئولیت مافوق یادآور شده‌اند که اگرچه دکترین سنتی بر رابطه مافوق با مادون^۲ استوار است، اما جوهره آن یعنی «تکلیف به نظارت»^۳ در مواجهه با سلاح‌های پیشرفته نیز زنده می‌ماند (Spadaro, 2023: 1125). بر این اساس، فرماندهی که بدون برقراری یک نظام نظارتی کارآمد، به خروجی یک الگوریتم اعتماد مطلق می‌کند، در واقع در انجام تکلیف پیشگیرانه خود قصور کرده است. این بازخوانی، «پیشگیری» را از یک مفهوم واکنشی میدانی به یک تکلیف معرفتی و فنی تبدیل می‌کند. در این پارادایم، فرمانده موظف است نه تنها از ارتکاب جرایم توسط زیردستان انسانی جلوگیری کند، بلکه باید از تبدیل شدن سامانه الگوریتمی به ابزار ارتکاب جنایت نیز پیشگیری نماید. این امر مستلزم آن است که فرمانده در صورت مشاهده هرگونه نقص در عملکرد سامانه یا تغییر در محیط عملیاتی که ریسک خطا را افزایش می‌دهد (مانند تغییر تراکم غیرنظامیان)، بلافاصله سامانه را تعلیق یا محدود کند. ناتوانی در انجام این تدابیر، یعنی ناتوانی در اعمال قدرت مادی برای پیشگیری، دقیقاً همان نقطه‌ای است که مسئولیت کیفری طبق ماده ۲۸ اساسنامه رم متبلور می‌شود. در نهایت، «پیشگیری» در قامت یک تکلیف نظارتی، تضمین می‌کند که خلأ مسئولیت ناشی از «شکاف کنترل» پر شود. با پیوند زدن پیشگیری به نظارت بر سیستم، حقوق کیفری بین‌المللی فرمانده را ملزم می‌کند که همواره «دستگیره قطع اضطراری»^۴ را چه به صورت فیزیکی و چه به صورت نظارتی در اختیار داشته باشد. از این منظر، هرگونه کوتاهی در برقراری یا استفاده از این سازوکارهای نظارتی، به مثابه «ترک فعل» مجرمانه در انجام وظایف فرماندهی تلقی خواهد شد.

۴- به سوی پارادایم مسئولیت‌پذیری انسانی

گذار از فرماندهی کلاسیک به جنگ الگوریتمی، صرفاً ابزار جنگ را تغییر نمی‌دهد؛ بلکه منطق انتساب مسئولیت را نیز تحت فشار قرار می‌دهد. اگر در الگوی سنتی، مسئولیت مافوق بر پیوند نسبتاً مستقیم میان «أمر انسانی» و «مرتکب انسانی» بنا شده بود، در محیطی که تصمیم‌گیری عملیاتی تا حدی به سامانه‌های خودمختار واگذار می‌شود، خطر اصلی نه فقط ارتکاب نقض، بلکه محو شدن نقطه اتکای مسئولیت است. از همین رو، گذار به پارادایم «مسئولیت‌پذیری انسانی» در واقع تلاشی است برای جلوگیری از تبدیل فناوری به سپری علیه پاسخگویی. در این پارادایم، انسان نباید به یک ناظر تشریفاتی یا امضاکننده صوری تقلیل یابد. مسئولیت‌پذیری انسانی مستلزم آن است که انسان همچنان در موقعیت تصمیم‌گیری معنادار، ارزیابی حقوقی، و امکان مداخله واقعی باقی بماند. کمیته بین‌المللی صلیب سرخ به درستی بر این نکته تأکید کرده است که نظام‌های سلاح خودمختار، به‌ویژه آنجا که بدون کنترل مؤثر انسانی به کار می‌روند، می‌توانند اجرای قواعد بنیادین حقوق بشردوستانه را با دشواری جدی مواجه کنند؛ از همین رو، این نهاد خواهان قواعد الزام‌آور جدید برای حفظ کنترل انسانی، محدودسازی دامنه عملکرد، و تضمین امکان نظارت و مداخله به موقع است (ICRC, 2021). این موضع، صرفاً یک ترجیح سیاست‌گذارانه نیست؛ بلکه بازتاب این واقعیت است که حقوق بین‌الملل بشردوستانه بدون یک سوژه انسانی پاسخگو، به تدریج قابلیت اعمال خود را از دست می‌دهد.

به لحاظ تحلیلی، مسئولیت‌پذیری انسانی سه لایه دارد. نخست، لایه پیشینی است؛ فرمانده باید پیش از استقرار یا استفاده از سامانه، از انطباق آن با قواعد حقوقی و محدودیت‌های عملیاتی اطمینان حاصل کند. اینجا ماده ۳۶ پروتکل الحاقی اول و منطق ارزیابی سلاح، اهمیت ویژه پیدا می‌کند. دوم، لایه حین عملیات است؛ انسان باید بتواند بر رفتار سامانه نظارت کند، سیگنال‌های

¹ - Meaningful Human Control

² - Superior-Subordinate

³ - Duty to oversee

⁴ - Kill Switch

خطر را بشناسد، و در صورت انحراف، مداخله مؤثر انجام دهد. سوم، لایه پسینی است؛ امکان حساب‌کشی، مستندسازی، و بازسازی تصمیم‌گیری باید وجود داشته باشد تا در صورت وقوع نقض، انتساب مسئولیت ممکن شود. اگر هر یک از این سه لایه حذف شود، «مسئولیت انسانی» به یک عنوان تهی فرو می‌کاهد. در اینجا باید میان «کنترل انسانی معنادار» و «حضور انسانی صوری» تمایز گذاشت. صرف آنکه یک انسان در زنجیره فرماندهی باقی بماند، برای تحقق مسئولیت‌پذیری کافی نیست. آنچه اهمیت دارد، واقعی بودن ظرفیت فهم، توقف، اصلاح، و امتناع از به‌کارگیری سامانه در شرایط پرخطر است. اگر فرمانده از نظر ساختاری نتواند خروجی الگوریتم را بفهمد یا از نظر زمانی فرصت مداخله نداشته باشد، وجود او در زنجیره تصمیم، نمی‌تواند مبنای معقولی برای رفع ابهام مسئولیت باشد. بنابراین، مسئولیت‌پذیری انسانی در معنای دقیق حقوقی آن، به معنای حفظ پیوند میان قدرت، آگاهی، و پاسخگویی است. از این منظر، پارادایم مسئولیت‌پذیری انسانی در برابر دو افراط ایستاده است؛ از یک‌سو، در برابر فناوری‌زدگی مطلق که اختیار را به سامانه می‌سپارد و از سوی دیگر، در برابر صوری‌سازی مسئولیت که انسان را فقط به عنوان «چهره حقوقی» نگه می‌دارد بی‌آنکه قدرت واقعی در اختیارش باشد. حقوق کیفری بین‌المللی اگر بخواهد در عصر الگوریتمی همچنان کارآمد بماند، باید بر این اصل پافشاری کند که هیچ معماری فنی نباید به حذف قابلیت انتساب کیفری منجر شود. به بیان دیگر، طراحی، استقرار، و به‌کارگیری سامانه‌های خودمختار باید به‌گونه‌ای باشد که مسئولیت انسانی نه فقط باقی بماند، بلکه قابل اثبات نیز باشد. نتیجه عملی این رویکرد آن است که فرماندهان، سیاست‌گذاران نظامی، و طراحان عملیات باید خود را نسبت به خروجی سامانه‌ها مسئول بدانند، حتی اگر سامانه بخشی از فرایند تصمیم را خودکار کرده باشد. این مسئولیت، البته، نامحدود نیست؛ اما محدودسازی آن نیز نمی‌تواند بر پایه ادعای «خودمختاری ماشین» صورت گیرد. در واقع، هرچه سطح خودمختاری بالاتر می‌رود، تکلیف انسانی برای نظارت، محدودسازی، و مداخله نیز باید سخت‌گیرانه‌تر شود. این همان نقطه‌ای است که حقوق بین‌الملل کیفری می‌تواند از سطح توصیف خطر به سطح هنجارسازی پاسخ برسد.

در جمع‌بندی، پارادایم مسئولیت‌پذیری انسانی را می‌توان ستون هنجاری بازخوانی ماده ۲۸ در عصر جنگ الگوریتمی دانست. این پارادایم نمی‌گذارد که «پیچیدگی فنی» به «خلأ مسئولیت» تبدیل شود. برعکس، آن را به تکلیف روشن تبدیل می‌کند؛ هرچا تصمیم‌گیری نظامی از انسان عبور می‌کند، باید همچنان به انسان بازگردد تا قابلیت نظارت، مداخله و پاسخگویی حفظ شود. بدون این بازگشت، نه فقط مسئولیت فرمانده، بلکه اعتبار خود رژیم حقوقی حاکم بر مداخلات مسلحانه تضعیف خواهد شد.

نتیجه‌گیری

قرائت دکتربین مسئولیت مافوق در عصر جنگ الگوریتمی، در نهایت، ما را به یک نقطه محوری می‌رساند: بحران مسئولیت کیفری فرماندهان نه در «خلأ هنجاری» بلکه در «محوشدگی عملیاتی» ارکان سنتی مسئولیت ریشه دارد. ماده ۲۸ اساسنامه رم، همچنان صحیح و الزام‌آور است؛ اما تفسیر متعارف آن که بر پیوندی نسبتاً مستقیم میان فرمانده، آگاهی قطعی از جرم، کنترل عملیاتی لحظه‌ای و امکان تحرک زمینی زیردست بنا شده، در مواجهه با سامانه‌های خودمختار که سرعت، پیچیدگی و عدم شفافیت محاسباتی را به صحنه مخاصمه وارد می‌کنند، از کفایت پیشینی خود فرو می‌ماند. بنابراین مسئله اصلی، جایگزینی ماده ۲۸ نیست؛ بلکه بازخوانی آن در پرتو واقعیت‌های تکنولوژیک است. این بازخوانی، چنان‌که در بخش‌های پیشین نشان داده شد، مستلزم تحول در هر یک از ارکان سه‌گانه است. نخست، «علم» نمی‌تواند به آگاهی قطعی و لحظه‌ای از جرم معین محدود بماند؛ در محیط الگوریتمی، این رکن به تکلیف فعال آگاهی از ریسک سیستمی تحول می‌یابد، بدان معنا که فرمانده نمی‌تواند از چشم‌پوشی شناختی نسبت به نقاط کور محاسباتی بهره ببرد. دوم، «کنترل مؤثر» از قدرت فرماندهی رو در روی زیردستان انسانی فاصله می‌گیرد و به کنترل بر معماری تصمیم‌گیری سامانه، یعنی بر مجموعه پارامترها، الگوریتم‌ها، محدوده‌های عملیاتی و سازوکارهای انسانی تعبیه‌شده در چرخه هدف‌گیری، تبدیل می‌شود. سوم، «پیشگیری» دیگر صرفاً صدور دستور توقف به یک عامل انسانی در آستانه ارتکاب نیست؛ بلکه به تکلیف نظارتی مستمر، شامل ارزیابی پیشینی سلاح، پایش برخط رفتار سامانه و تضمین وجود سازوکار مداخله و توقف اضطراری، بازتعریف می‌شود. در هر سه رکن، منطق واحد حاکم است: هیچ‌یک از این وجوه نباید اجازه دهد که «پیچیدگی فنی» به «رهایی از پاسخگویی» بدل شود. از همین جاست که مفهوم «مسئولیت‌پذیری انسانی» به‌عنوان ستون

هنجاری این بازخوانی وارد می‌شود. مسئولیت‌پذیری انسانی در معنای دقیق، نه یک شعار اخلاقی، بلکه یک معیار فنی-حقوقی برای طراحی و به‌کارگیری سامانه‌هاست: قابلیت درک منطق تصمیم، قابلیت مداخله به‌موقع، قابلیت توقف اضطراری، و قابلیت بازسازی حساب‌پذیر تصمیم‌گیری. این معیار در نقطه تقاطع سه جریان هنجاری‌پیشین مقاله قرار می‌گیرد: موضع صریح کمیته بین‌المللی صلیب سرخ مبنی بر ضرورت حفظ کنترل و نظارت انسانی و درخواست قواعد الزام‌آور جدید، تکلیف ارزیابی قانونی سلاح‌های جدید در ماده ۳۶ پروتکل الحاقی اول، و منطق مسئولیت‌ناظر بر کوتاهی مافوق در رویه قضایی محاکم بین‌المللی (به‌ویژه پرونده اوریچ و بمبا) در این چارچوب، «کنترل انسانی» دیگر یک عنوان فرمایشی نیست؛ بلکه پیش‌شرط امکان‌پذیر بودن خود مسئولیت کیفری است.

بر این اساس، می‌توان مدل پیشنهادی زیر را برای تفسیر ماده ۲۸ در عصر جنگ الگوریتمی ارائه کرد. این مدل بر نفی هرگونه «خودمختاری سپرمانند» استوار است و سه قاعده هنجاری را پیشنهاد می‌کند. قاعده نخست (اصل نسبیت خودمختاری): سطح خودمختاری سامانه، نه کاهش‌دهنده و نه حذف‌کننده مسئولیت فرمانده، بلکه تعیین‌کننده کانون تکلیف اوست؛ هرچه خودمختاری بیشتر شود، تکلیف به طراحی پیشینی، نظارت برخط و مداخله مؤثر سنگین‌تر می‌شود. قاعده دوم (اصل تقدم کنترل انسانی معنادار): هیچ سامانه‌ای را نمی‌توان به‌کار گرفت که ساختاراً امکان ادراک حقوقی، توقف اضطراری و امتناع آگاهانه را از انسان سلب کند؛ فقدان این قابلیت‌ها به معنای عدم تحقق «کنترل مؤثر» و باقی‌ماندن بار مسئولیت بر فرمانده است. قاعده سوم (اصل نسبیت زدایی از دفاع تکنولوژیک): فرمانده نمی‌تواند صرف «خطای ماشین» یا «پیچیدگی ناشفاف الگوریتم» را به‌عنوان دلیل رفع مسئولیت مطرح کند؛ مگر آن‌که اثبات کند در انجام تکلیف فعال آگاهی، نظارت و پیشگیری مطابق با مراتب بالای مراقبت قابل انتظار از یک فرمانده در وضعیت عینی مشابه قصور نکرده است. اجرای این مدل، البته، متضمن پاسخ به یک پرسش هنجاری عمیق‌تر است: آیا حقوق کیفری بین‌المللی باید «خودمختاری ماشین» را یک واقعیت فنی بپذیرد و صرفاً واکنش نشان دهد، یا باید به‌عنوان یک هنجار تنظیمی، «خودمختاری قابل قبول» را از پیش محدود کند؟ استدلال این مقاله بر آن است که حقوق کیفری بین‌المللی نمی‌تواند منفعل بماند. اگر این رژیم حقوقی در اواخر قرن بیستم برای «مرگ صدها هزار نفر» در مخاصمات، اشخاص حقیقی را پاسخگو دانست، در قرن بیست‌ویکم نیز نمی‌تواند با تفویض تصمیم‌های مرگ و زندگی به معماری‌های محاسباتی، نقطه اتکای مسئولیت را محو کند. تداوم حیات حقوق کیفری بین‌المللی، نه در توانایی آن در مهار کامل فناوری، بلکه در مصممیت آن در باقی‌نگه‌داشتن انسان به‌عنوان تنها سوژه‌ای است که می‌تواند بمیرد، و تنها سوژه‌ای است که می‌تواند پاسخ‌گو باشد.

به همین دلیل، دستاورد این بازخوانی را می‌توان چنین خلاصه کرد؛ از یک‌سو، ماده ۲۸ را از تَصَلُّب تفسیری تطبیق‌نیافتنی با محیط الگوریتمی رها می‌کند، و از سوی دیگر، به آن حیاتی تازه می‌بخشد تا همچنان کارکرد بنیادین خود را ایفا کند یعنی تضمین اینکه «فرمانده بودن» هرگز به «بی‌مسئولیت بودن» نیانجامد. درست در همین نقطه است که مسئولیت فرمانده در عصر جنگ الگوریتمی، نه به‌مثابه یک نهاد رو به افول، بلکه به‌مثابه چهارراه حیاتی حقوق بشردوستانه، حقوق کیفری بین‌المللی و در نهایت، خود اخلاق مسئولیت، ظاهر می‌شود.

1. Abraham, Y. (2024, April 3). “‘Lavender’: The AI machine directing Israel’s bombing spree in Gaza”. +972 Magazine. <https://www.972mag.com/lavender-ai-israeli-army-gaza/>
2. Ambos, K. (2014). *Treatise on International Criminal Law: Volume I: Foundations and General Part*. Oxford University Press.
3. Andersin, E. (2025). The use of the ‘Lavender’ in Gaza and the law of targeting. *Journal of International Humanitarian Legal Studies*, 16(2), 336-360. <https://doi.org/10.1163/18781527-bja10074>
4. Bátori, A. (2025). International criminal law issues of autonomous weapons: The Kargu-2 drone in Libya. *Gradus*, 12(1), 1-9. <https://doi.org/10.47833/2025.1.ART.011>
5. Blank, L. R. (2024, June 28). “The Gospel, Lavender, and the law of armed conflict”. Lieber Institute, West Point. <https://lieber.westpoint.edu/gospel-lavender-law-armed-conflict/>
6. Breaking Defense. (2024). “Governing AI under fire in Ukraine”. The Cairo Review of Global Affairs. <https://www.thecairoreview.com/essays/governing-ai-under-fire-in-ukraine/>
7. Cassese, A. (2013). *Cassese’s International Criminal Law* (third Ed.). Oxford University Press.
8. Chengeta, T. (2016). Accountability gap: Autonomous weapon systems and modes of responsibility in international law. *Denver Journal of International Law and Policy*, 45(1), 1-50. <https://digitalcommons.du.edu/djilp/vol45/iss1/3/>
9. Crawford, Emily. Pert, Alison. (2024). *International Humanitarian Law*. Third Ed. Cambridge: Cambridge University Press.
10. Goodwin, A. (2024, April 3). “Israel used AI to identify targets in Gaza, sources claim”. The Guardian. <https://www.theguardian.com/world/2024/apr/03/israel-gaza-ai-database-hamas-airstrikes>
11. Gunawan, Y. et al. (2022). Command responsibility of autonomous weapons under international humanitarian law. *Cogent Social Sciences*, 8(1), 1-16. <https://doi.org/10.1080/23311886.2022.2139906>
12. Horowitz, M. C., & Scharre, P. (2023). “AI and international stability: Risks and confidence-building measures”. Center for a New American Security. <https://www.cnas.org/publications/reports/ai-and-international-stability-risks-and-confidence-building-measures>
13. ICRC. (2018). “Ethics and autonomous weapon systems: An ethical basis for human control? International Committee of the Red Cross”. <https://www.icrc.org/en/document/ethics-and-autonomous-weapon-systems-ethical-basis-human-control>
14. International Committee of the Red Cross (ICRC). (2019). Artificial intelligence and machine learning in armed conflict: A human-centred approach. ICRC. Retrieved from <https://www.icrc.org/en/document/artificial-intelligence-and-machine-learning-armed-conflict-human-centred-approach>
15. International Committee of the Red Cross. (2021, May 12). ICRC position on autonomous weapon systems: ICRC position and background paper. International Review of the Red Cross. <http://international-review.icrc.org/articles/icrc-position-on-autonomous-weapon-systems-icrc-position-and-background-paper-915>
16. International Criminal Court (ICC). (2016, March 21). Prosecutor v. Jean-Pierre Bemba Gombo [Judgment pursuant to Article 74 of the Statute]. Case No. ICC-01/05-01/08.
17. International Criminal Court (ICC). (2016, March 21). Prosecutor v. Jean-Pierre Bemba Gombo [Judgment pursuant to Article 74 of the Statute]. Case No. ICC-01/05-01/08. Retrieved from <https://www.icc-cpi.int/>
18. International Criminal Court (ICC). (2018, June 8). Prosecutor v. Jean-Pierre Bemba Gombo [Judgment on the appeal of Mr. Jean-Pierre Bemba Gombo against Trial Chamber III’s Judgment]. Case No. ICC-01/05-01/08 A.
19. International Criminal Court (ICC). (2018, June 8). Prosecutor v. Jean-Pierre Bemba Gombo [Judgment on the appeal of Mr. Jean-Pierre Bemba Gombo against Trial Chamber III’s Judgment]. Case No. ICC-01/05-01/08 A. Retrieved from <https://www.icc-cpi.int/>
20. International Criminal Tribunal for the former Yugoslavia (ICTY). (1998, November 16). Prosecutor v. Zejnil Delalić, Zdravko Mucić (aka “Pavo”), Hazim Delić and Esad Landžo (aka “Zenga”) [Judgment]. Case No. IT-96-21-T (the Čelebići case).
21. International Criminal Tribunal for the former Yugoslavia (ICTY). (2008, April 22). Prosecutor v. Enver Hadžihasanović and Amir Kubura [Judgment]. Case No. IT-01-47-A.
22. International Criminal Tribunal for the former Yugoslavia (ICTY). (1998, November 16). Prosecutor v. Zejnil Delalić, Zdravko Mucić (aka “Pavo”), Hazim Delić and Esad Landžo (aka “Zenga”) [Judgment]. Case No. IT-96-21-T (the Čelebići case). Retrieved from <https://www.icty.org/>
23. International Criminal Tribunal for the former Yugoslavia (ICTY). (2008, April 22). Prosecutor v. Enver Hadžihasanović and Amir Kubura [Judgment]. Case No. IT-01-47-A. Retrieved from <https://www.icty.org/>
24. Kania, E. B. (2017). “Battlefield singularity: Artificial intelligence, military revolution, and China’s future military power”. Center for a New American Security. <https://www.cnas.org/publications/reports/battlefield-singularity-artificial-intelligence-military-revolution-and-chinas-future-military-power>

25. Liu, H.-Y. (2012). Categorization and legality of autonomous and remote weapons systems. *International Review of the Red Cross*, 94(886), 627-652. <https://doi.org/10.1017/S181638311300012X>
26. Payne, K. (2021). *I, warbot: The dawn of artificially intelligent conflict*. Hurst.
27. Rome Statute of the International Criminal Court. (1998). *United Nations, Treaty Series*, vol. 2187, No. 38544.
28. Schabas, W. A. (2016). *The International Criminal Court: A Commentary on the Rome Statute* (second Ed.). Oxford University Press.
29. Scharre, P. (2018). *Army of none: Autonomous weapons and the future of war*. W. W. Norton & Company.
30. Schmitt, M. N., & Thurnher, J. S. (2013). "Out of the loop": Autonomous weapon systems and the law of armed conflict. *Harvard National Security Journal*, 4, 231-281. <https://harvardnsj.org/2013/05/22/out-of-the-loop-autonomous-weapon-systems-and-the-law-of-armed-conflict/>
31. Sehwat, V. (2020). Autonomous Weapon System and Command Responsibility. *Florida Journal of International Law*. 31(3), 315-337. <https://scholarship.law.ufl.edu/fjil/vol31/iss3/2>
32. SIPRI. (2024). *The war in Ukraine and the role of autonomous systems in modern warfare*. Stockholm International Peace Research Institute.
33. Spadaro, A. (2023). A Weapon is No Subordinate: Autonomous Weapon Systems and the Scope of Superior Responsibility. *Journal of International Criminal Justice*, 21(5), 1119-1136. <https://doi.org/10.1093/ijcj/mqad025>
34. United Nations Security Council (UNSC). (2021, March 8). Final report of the Panel of Experts on Libya established pursuant to Resolution 1973 (2011) (S/2021/229). Retrieved from <https://digitallibrary.un.org/>
35. United Nations Security Council. (2021, March 8). Letter dated 8 March 2021 from the Panel of Experts on Libya established pursuant to resolution 1973 (2011) (S/2021/229). <https://digitallibrary.un.org/record/3905159>
36. United States Military Commission. (1946). Trial of General Tomoyuki Yamashita. Law Reports of Trials of War Criminals, Vol. IV. London: HMSO.

Criminal Responsibility of Commanders in the Age of Algorithmic Warfare; Rereading the Criteria of "Science," "Effective Control," and "Prevention" in the Light of International Criminal Law

Pouria Ahmadi

graduate in Criminal Law and Criminology from the Islamic Azad University, Semnan Branch, Iran

Amin Alizadeh (Corresponding Author)

PhD Candidate, Department of Criminal Law and Criminology, University of Tehran (Campus), Iran.

Abstract

Algorithmic warfare has brought about a fundamental transformation in the place of decision-making on the battlefield; a decision that was once formed in the mind and will of a human commander is now increasingly produced within automated systems and through unpredictable computational processes. This shift in the "center of decision" raises a legal question, the neglect of which jeopardizes the meaning of criminal responsibility in the international arena; in a system where the commander no longer issues explicit orders but rather defines the frameworks within which a system makes independent decisions, what is the role of attributing responsibility to him? The present article, adopting a descriptive-analytical method, seeks to answer the question of whether the framework of criminal responsibility of commanders in Article 28 of the Rome Statute can maintain its interpretative capacity in the face of the challenges of algorithmic warfare or whether it collapses in the face of the reality of autonomous systems. The main finding of the research is that the "control break" resulting from the autonomy of systems does not make Article 28 obsolete, but calls for a serious dialectic with new conditions; but this dialectic requires a critical rereading of three fundamental criteria of responsibility: the transition from "definitive science" to "risk-based science", the reformulation of "effective control" as "systemic control", and the redefinition of "prevention" as a task of continuous monitoring and meaningful human intervention. The article ultimately proposes an analytical model that, by preserving the text of Article 28 and without an arbitrary expansion of criminal responsibility, keeps human agency alive at the heart of international criminal law, while providing international judges and prosecutors with practical guidance for establishing responsibility in algorithmic battlefields.

Keywords: algorithmic warfare, criminal responsibility of commanders, Article 28 of the Rome Statute, effective control, meaningful human control, human agency.