

УДК 004.8

***ЭТИЧЕСКИЕ, ПРАВОВЫЕ И СТРАТЕГИЧЕСКИЕ АСПЕКТЫ
ПРИМЕНЕНИЯ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА В ВОЕННОЙ
СФЕРЕ***

Бойкова А. В.

д.э.н., доцент,

ФГБОУ ВО «Тверской государственный технический университет»,

Тверь, Россия

Рыльцин И.А.

к.т.н., докторант

ФГК ВОУ ВО "Военная академия воздушно-космической обороны имени

маршала Советского Союза Г.К. Жукова" Министерства обороны Российской

Федерации,

Тверь, Россия

Павленко В. М.

разработчик

АО «Северсталь-инфоком»

Москва, Россия

Аннотация. Статья посвящена исследованию этических, правовых и стратегических вызовов, возникающих при интеграции технологий искусственного интеллекта в военную сферу. В работе рассмотрены ключевые этические проблемы, связанные с применением летальных автономных систем, включая вопрос о допустимости передачи машине решения о лишении жизни человека и проблему «пробелов ответственности». Проанализировано современное состояние международно-правового регулирования военного ИИ, выявлены причины отсутствия обязывающего договора и роль Политической декларации об ответственном военном ИИ. Исследованы стратегические риски, включая снижение порога войны, гонку вооружений в сфере ИИ и угрозу

вовлечения негосударственных акторов. Сделан вывод о необходимости выработки глобальных норм и механизмов контроля, обеспечивающих баланс между инновациями и безопасностью..

Ключевые слова: искусственный интеллект, военная сфера, летальные автономные системы, международное гуманитарное право, этика, стратегическая стабильность, регулирование.

***ETHICAL, LEGAL AND STRATEGIC ASPECTS OF ARTIFICIAL
INTELLIGENCE APPLICATION IN THE MILITARY DOMAIN***

Boykova A. V.

PhD, Associate Professor,

Tver State Technical University,

Tver, Russia

Ryltsin I.A.

PhD, Doctoral Candidate

Federal State Budgetary Educational Institution of Higher Education "Marshal of the Soviet Union G.K. Zhukov Military Academy of Aerospace Defense" of the Ministry of Defense of the Russian Federation,

Tver, Russia

Pavlenko V. M.

developer

JSC Severstal-infocom

Moscow, Russia

Abstract. The article is devoted to the study of ethical, legal and strategic challenges arising from the integration of artificial intelligence technologies into the military domain. The paper examines key ethical issues associated with the use of lethal autonomous systems, including the question of whether it is permissible to delegate the decision to take a human life to a machine and the problem of "responsibility gaps." The current state of international legal regulation of military AI is analyzed,

the reasons for the absence of a binding treaty and the role of the Political Declaration on Responsible Military AI are identified. Strategic risks are investigated, including the lowering of the threshold for war, the arms race in AI, and the threat of involvement of non-state actors. It is concluded that it is necessary to develop global norms and control mechanisms that ensure a balance between innovation and security.

Keywords: artificial intelligence, military domain, lethal autonomous systems, international humanitarian law, ethics, strategic stability, regulation.

Интеграция технологий искусственного интеллекта (ИИ) в вооружения и военную организацию представляет собой одно из наиболее значимых направлений современной технологической трансформации оборонных систем. Вместе с тем, внедрение ИИ в военную сферу порождает множество сложных вопросов этического, юридического и стратегического характера, требующих комплексного анализа и выработки адекватных механизмов регулирования.

Цель настоящей работы – систематизация и анализ основных вызовов, связанных с применением ИИ в военной области, включая этические дилеммы, правовые пробелы и стратегические риски, а также определение возможных направлений формирования международно-правового режима регулирования военного ИИ.

Главный этический вопрос, возникающий в контексте применения военных ИИ-систем, – допустимо ли передавать машине решение о лишении жизни человека. Многие философы, правозащитники и религиозные деятели считают это морально неприемлемым, утверждая, что лишение жизни должно всегда оставаться под ответственным контролем человека. Появление летальных автономных систем (LAWS), где алгоритм выбирает цель и открывает огонь без санкции оператора, называют переходом «красной линии». Генеральный секретарь ООН Антониу Гутерреш неоднократно заявлял, что оружие, которое самостоятельно выбирает и поражает цель, по политическим и

Дневник науки | www.dnevniknauki.ru | СМИ Эл № ФС 77-68405 ISSN 2541-8327

моральным соображениям должно быть запрещено. Многие страны, включая почти все государства Европейского Союза, поддерживают эту точку зрения и выступают за запрет «убийственных роботов» [1].

Не в полной мере этична и сама мысль, что алгоритм, который не обладает ни пониманием ценности человеческой жизни, ни способностью к сочувствию, будет принимать столь судьбоносные решения. Традиционные нормы военного права не знают понятия «вина машины», соответственно, ответственность ляжет на операторов или командование. Однако могут возникать так называемые «пробелы ответственности» (responsibility gaps). Чтобы не допустить размытия ответственности, потребуется устанавливать строгие протоколы использования ИИ и, возможно, вводить новую юридическую категорию для таких случаев.

Кроме того, в моделях ИИ могут содержаться предубеждения разработчиков. Например, алгоритмы распознавания лиц нередко хуже различают людей определённой этнической группы. Это бросает вызов принципу справедливости и гуманности войны. Этические руководства, подобные тем, что приняты в Министерстве обороны США и НАТО, призваны минимизировать такие эффекты, однако воплотить их в код непросто [2].

Международное гуманитарное право (МГП) строится на следующих фундаментальных принципах: различие (нельзя целить в гражданских лиц), пропорциональность (ущерб гражданским лицам не должен чрезмерно превышать военную необходимость) и милосердие (запрет излишних страданий). Ключевой вопрос заключается в том, может ли ИИ гарантированно соблюдать эти принципы.

На глобальном уровне на сегодняшний день не существует обязывающего договора, регулирующего военные ИИ-системы. С 2014 года в рамках Конвенции ООН о некоторых видах обычного оружия (CCW) проходят встречи экспертов по летальным автономным системам. Однако прогресс остаётся медленным — консенсус между государствами-участниками

Дневник науки | www.dnevniknauki.ru | СМН ЭЛ № ФС 77-68405 ISSN 2541-8327

отсутствует. В 2021 году Генеральный секретарь ООН впервые официально призвал начать переговоры о договоре к 2026 году. В 2022–2023 годах более 60 стран (включая государства ЕС и развивающиеся страны) поддержали идею юридически обязывающего инструмента. Однако против выступают ключевые военные державы – США, Россия, Китай, Индия, которые предпочитают формулировку о «достаточности существующего права» и национальных руководствах. Это блокирует заключение нового договора, поскольку в рамках ООН обычно требуется консенсус всех крупных игроков [3].

Альтернативой стала Политическая декларация об ответственном военном ИИ, инициированная США. Она не имеет силы закона, но фиксирует ряд норм: ИИ в вооружённых конфликтах должен применяться в соответствии с международным правом; государства обязуются принимать меры против рисков некорректной работы ИИ; сохранять человеческую ответственность. К концу 2024 года декларацию подписали около 60 государств, однако она не включает ряд ключевых стран (РФ, КНР и другие не присоединились). Таким образом, правовой вакуум в значительной степени сохраняется, что представляет серьёзный вызов: технология развивается быстрее права [4].

Отдельный аспект – применимость экспортного контроля и режимов нераспространения. Проводятся параллели с ядерным оружием: предлагается создать «ИИ-нераспространительный режим» по аналогии с ядерным. Однако ИИ – технология двойного назначения, широко распространённая и, в отличие от обогащённого урана, крайне диффузная. Невозможно контролировать алгоритмы и программное обеспечение так же, как ядерные материалы. США предпринимают попытки точечно ограничивать экспорт передовых чипов в Китай, однако эффективность таких мер вызывает сомнения, и они провоцируют встречные шаги. Таким образом, классические механизмы контроля вооружений с трудом применимы к сфере ИИ. Некоторое время обсуждалась идея создания глобального органа по надзору за ИИ – аналога

МАГАТЭ, однако консенсус отсутствует, и сложно представить инспекторов, проверяющих военные алгоритмы, в силу их высокой секретности.

Военный ИИ несёт риски дестабилизации стратегической обстановки по нескольким направлениям.

Во-первых, снижение порога войны: если часть задач (разведка, удары) выполняют машины, потери личного состава уменьшаются, что может побудить политиков легче решиться на применение силы, полагаясь на «бескровные» технологии. Это может привести к более частым вооружённым инцидентам.

Во-вторых, ускорение войны грозит тем, что решения о жизни и смерти могут приниматься слишком быстро для дипломатического урегулирования. Концепция «hyperwar» (гипервойны) описывает конфликт, где ИИ по обе стороны обмениваются ударами на сверхчеловеческих скоростях, что практически исключает возможность деэскалации или корректировки курса, – это крайне опасная перспектива.

Гонка военных ИИ между великими державами уже идёт. Она может привести к тому, что стороны будут принимать на вооружение сырой, недостаточно протестированный ИИ из опасения отстать – так называемый «синдром спешки». Это напоминает ситуацию с ядерным оружием на ранних этапах, когда также испытывали системы поспешно, пока не осознали масштаб угрозы. Однако если в случае ядерного оружия осознание пришло сравнительно быстро (через 15 лет после 1945 года уже появились договоры), то здесь понимание последствий размыто по множеству мелких инцидентов, а не сконцентрировано в одном событии [5].

Вовлечение негосударственных акторов представляет ещё одну значимую угрозу. ИИ-технологии доступны коммерчески, поэтому террористические или преступные группировки также могут их применять. Уже были попытки оснастить самодельные дроны примитивными системами ИИ для наведения. В перспективе террорист может запустить рой микродронов с системами

Дневник науки | www.dnevniknauki.ru | СМИ ЭЛ № ФС 77-68405 ISSN 2541-8327

распознавания лиц для атаки на конкретных людей – это новая форма угрозы, к которой международное сообщество не готово. Здесь требуется международный контроль и сотрудничество правоохранительных органов.

Доверие и психологический фактор также играют важную роль. Военным командирам необходимо время, чтобы начать доверять ИИ-системам. Известны случаи, когда операторы отключали автоматические режимы из опасения ошибки (например, во время войны в Персидском заливе система Patriot ошибочно опознала свой самолёт как вражеский – после этого многие операторы стали с осторожностью относиться к её сигналам). В перспективе, если ИИ докажет свою надёжность, могут возникнуть обратные перекосы – чрезмерное доверие к алгоритму, когда человек перестаёт критически оценивать решения машины и становится просто исполнителем. Оптимальным представляется разумное сотрудничество человека и ИИ, однако его достижение требует подготовки кадров, новой военной педагогики и изменения организационной культуры.

Наконец, этическое лидерство и имиджевые издержки. Если одна сторона будет широко использовать ИИ-системы для поражения целей, она может потерять моральное превосходство и поддержку мирового сообщества. Появление видео, демонстрирующих, как беспилотный робот расстреливает убегающих солдат, может вызвать общественное отвращение. Поэтому даже выиграв битву, можно проиграть информационную войну. Страны НАТО осознают это и потому настаивают на ответственном, прозрачном подходе, чтобы сохранить легитимность своих действий.

Проведённый анализ позволяет сделать вывод, что военный ИИ представляет собой «палку о двух концах». С одной стороны, он обещает повышение эффективности и спасение жизней своих солдат, с другой – несёт риски для жизни гражданского населения и стабильности мира в целом. Как отмечают эксперты Carnegie Endowment, отсутствие глобальных рамок для

военного ИИ создаёт опасный регулирующий вакуум, чреватый угрозой миру и безопасности.

Чтобы ИИ стал больше благом, чем злом, необходимо срочное сотрудничество международного сообщества по следующим направлениям: выработка политических норм и правил применения военного ИИ; обмен опытом по безопасному внедрению; создание механизмов контроля эскалации (например, прямых линий связи на случай инцидентов с автономными системами). Параллельно важно продвигать исследования по безопасности ИИ (AI safety) применительно к военной области, вкладывать средства в методы объяснимости ИИ, защиты от взлома и этического проектирования систем.

Как справедливо отмечено в резолюции Европейского парламента, «ИИ изменит войну, но не должен изменить наши ценности; человек должен оставаться в центре решений, а международное право – распространяться на все новые технологии». Следующие несколько лет станут решающими – удастся ли человечеству направить военный ИИ в русло сдержанности и законности или мы вступим в эру непросчитанных рисков. Только осознанный подход, совмещающий инновации с ответственностью, позволит воспользоваться преимуществами ИИ, избежав техногенных угроз, о которых предупреждают футурологи.

Библиографический список:

1. 'Politically unacceptable, morally repugnant': UN chief calls for global ban on 'killer robots' [Электронный ресурс]. Режим доступа: URL: <https://news.un.org/en/story/2025/05/1163256> (дата обращения: 01.07.2026).

2. Rosen B. How to Make Military AI Governance More Robust [Электронный ресурс]. Режим доступа: URL: <https://warontherocks.com/how-to-make-military-ai-governance-more-robust/> (дата обращения: 01.07.2026).

3. NATO and artificial intelligence: navigating the challenges and opportunities [Электронный ресурс]. Режим доступа: URL: <https://www.nato-dnevnik-nauki.ru> | www.dnevniknauki.ru | СМИ ЭЛ № ФС 77-68405 ISSN 2541-8327

pa.int/document/2024-nato-and-ai-report-clement-058-stc (дата обращения: 01.07.2026).

4. Carnegie Endowment for International Peace. Governing Military AI Amid a Geopolitical Minefield [Электронный ресурс]. Режим доступа: URL: <https://carnegieendowment.org/research/2024/07/governing-military-ai-amid-a-geopolitical-minefield> (дата обращения: 01.07.2026).

5. European Parliament Research Service. Artificial Intelligence in Warfare: Ethical and Legal Challenges [Электронный ресурс]. Режим доступа: URL: [https://www.europarl.europa.eu/RegData/etudes/BRIE/2025/769580/EPRS_BRI\(2025\)769580_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2025/769580/EPRS_BRI(2025)769580_EN.pdf) (дата обращения: 01.07.2026).