

УДК 004.056.5

Модель угроз для систем искусственного интеллекта в киберпространстве: вызовы 2025 года

Авхадиев Д.И.

Независимый исследователь в области IT и кибербезопасности

E-mail: donat.avkh@gmail.com

Искусственный интеллект совершает революцию в сфере кибербезопасности, дополняя традиционные механизмы защиты передовыми алгоритмами и средствами автоматизации. Искусственный интеллект, машинное обучение и глубокое обучение играют ключевую роль в автоматизации важнейших процессов кибербезопасности. Используя эти технологии, организации могут оптимизировать такие задачи, как обнаружение угроз, реагирование на инциденты и оценка уязвимостей. Роль генеративного искусственного интеллекта в сфере кибербезопасности становится всё более значимой в современную цифровую эпоху. Искусственный интеллект анализирует огромные объёмы данных в режиме реального времени, выявляет потенциальные угрозы и прогнозирует будущие атаки. Он совершенствует технологии, которые компании используют для борьбы с киберпреступниками, и помогает организациям обеспечивать безопасность данных клиентов. Целью настоящей статьи является анализ современных вызовов в области применения искусственного интеллекта в киберпространстве. В работе рассмотрены угрозы, направленные на модели искусственного интеллекта, такие как состязательные атаки, отправление обучающих выборок и манипуляции алгоритмами. Также затронуты вопросы этики, прозрачности и регулирования. Представленные результаты подчеркивают необходимость комплексной оценки рисков при интеграции искусственного интеллекта в системы цифровой безопасности.

Ключевые слова: искусственный интеллект, кибербезопасность, киберпреступления, защита, киберпространство, машинное обучение, кибератака.

Threat Model for Artificial Intelligence Systems in Cyberspace: Challenges of 2025

Donat I. Avkhadiev

Independent Researcher in IT and Cybersecurity

E-mail: donat.avkh@gmail.com

Abstract

Artificial intelligence is revolutionizing cybersecurity by complementing traditional protection mechanisms with advanced algorithms and automation tools. Artificial intelligence, machine learning, and deep learning play a key role in automating critical cybersecurity processes. Using these technologies, organizations can optimize tasks such as threat detection, incident response, and vulnerability assessment. The role of generative artificial intelligence in the field of cybersecurity is becoming increasingly important in the modern digital age. Artificial intelligence analyzes huge amounts of data in real time, identifies potential threats and predicts future attacks. He improves the technologies that companies use to fight cybercriminals and helps organizations ensure the security of customer data. The purpose of this article is to analyze modern challenges in the field of artificial intelligence in cyberspace. The paper considers threats aimed at artificial intelligence models, such as adversarial attacks, sending training samples and manipulating algorithms. Issues of ethics, transparency and regulation were also raised. The presented results emphasize the need for a comprehensive risk assessment when integrating artificial intelligence into digital security systems.

Keywords: artificial intelligence, cybersecurity, cybercrime, protection, cyberspace, machine learning, cyberattack.

Киберугрозы представляют собой преднамеренные действия, осуществляемые злоумышленниками с целью нарушения работы, нанесения ущерба или дестабилизации процессов в цифровом пространстве [3]. В противовес этому, кибербезопасность охватывает комплекс превентивных мер, направленных на защиту цифровых ресурсов, включая информацию, данные и программные компоненты, от всевозможных атак. С развитием технологий злоумышленники демонстрируют всё большую изобретательность и способность обходить традиционные механизмы защиты. Их опережающие действия усложняют обеспечение надёжной безопасности как для организаций, так и для отдельных пользователей.

Согласно последнему отчёту Verified Market Research, объём глобального рынка ИИ в области кибербезопасности достиг 7,58 миллиарда долларов в 2024 году, а к 2030 году прогнозируется его увеличение до 80,83 миллиарда долларов [12]. В условиях стремительного роста числа кибератак это не является неожиданностью, учитывая, что злоумышленники также используют современные технологические инструменты.

Машинное обучение существенно повлияло на развитие бизнеса, позволив компаниям принимать более обоснованные решения, автоматизировать рутинные процессы и извлекать ценные аналитические сведения. Тем не менее, интеграция генеративного ИИ в кибербезопасностные системы сопровождается рядом вызовов и потенциальных угроз.

Цель исследования

Основные проблемы, связанные с внедрением ИИ в сферу кибербезопасности [5]:

- Манипулирование данными: системы искусственного интеллекта используют данные для обучения и выявления закономерностей. Хакеры могут получить доступ к этим обучающим данным и манипулировать ими, внося искажения или изменяя их в своих интересах. Такие действия могут снизить точность и эффективность моделей ИИ и потенциально привести к нарушениям безопасности.

- Кибератаки с использованием искусственного интеллекта: хакеры могут применять методы искусственного интеллекта для разработки интеллектуальных вредоносных программ, которые могут адаптироваться и модифицироваться, чтобы избежать обнаружения даже самым современным программным обеспечением для кибербезопасности. Из-за этого традиционным мерам безопасности становится всё сложнее противостоять развивающимся угрозам.

- Доступность данных: эффективность моделей ИИ во многом зависит от количества и качества данных, доступных для обучения. Если обучающих данных недостаточно или они

необъективны, система ИИ может работать не так, как ожидалось. Недостаточно обученные модели могут выдавать ложные срабатывания, создавая ложное ощущение безопасности и делая организации уязвимыми для незамеченных угроз.

- Проблемы конфиденциальности: для точного понимания моделей поведения пользователей моделям ИИ часто требуется доступ к реальным пользовательским данным [5]. Без соответствующих мер, таких как маскировка или шифрование данных, эти конфиденциальные данные становятся уязвимыми для нарушений конфиденциальности, что может подвергнуть пользователей риску и сделать их жертвами злоумышленников.

- Атаки на системы искусственного интеллекта: сами системы искусственного интеллекта не защищены от кибератак. Хакеры могут намеренно подменять или манипулировать данными, поступающими в эти модели, изменяя их поведение в соответствии со своими злонамеренными целями. Это подчеркивает важность защиты систем искусственного интеллекта от злонамеренных манипуляций.

Результаты исследования

В условиях стремительного развития технологий обеспечение защиты сетевой инфраструктуры организаций приобретает решающее значение, при этом управление уязвимостями выступает в качестве ключевого компонента этой стратегии. Современные компании ежедневно сталкиваются с множеством угроз, которые необходимо своевременно идентифицировать, классифицировать и нейтрализовать, чтобы избежать возможного ущерба. В этом контексте значительную роль играют достижения в области искусственного интеллекта, предлагающие новые подходы к анализу и оценке эффективности мер информационной безопасности.

С усовершенствованием ИИ-систем их способности к обучению в реальных сценариях существенно усложняют действия злоумышленников, стремящихся обойти защиту. Принимая во внимание объём трафика, ежедневно циркулирующего между серверами и клиентскими устройствами, задача своевременного обнаружения потенциальных инцидентов становится всё более трудоёмкой для специалистов по кибербезопасности [2].

В цифровую эпоху, на фоне технологического прогресса, искусственный интеллект демонстрирует значительные успехи, параллельно порождая комплекс этических вызовов, требующих всестороннего анализа. По мере внедрения ИИ в такие сферы, как автономный транспорт, интеллектуальные устройства, медицинское обслуживание и образование, возрастает необходимость в осмысленном подходе к оценке его социального воздействия.

Один из ключевых факторов, повлиявших на эволюцию искусственного интеллекта, — это расширение возможностей вычислительных ресурсов. За последние годы

производительность и скорость современных вычислительных систем значительно возросли, что открыло исследователям путь к обучению сложных моделей машинного обучения и обработке больших массивов данных в режиме реального времени. Рост вычислительной мощности сопровождался усовершенствованием аппаратных решений, в частности, появлением высокопроизводительных процессоров и графических ускорителей (GPU) [9].

Наряду с этим, развитие программного обеспечения и прогресс в области архитектуры обработки информации способствовали созданию более эффективных алгоритмов и технологий параллельных вычислений. Существенным фактором также является широкая доступность облачных платформ, обеспечивающих аренду вычислительных мощностей по мере необходимости, что делает ИИ-технологии более масштабируемыми и гибкими в применении.

Параллельно с этим, распространение Интернета вещей стимулирует рост объема данных, поступающих от подключённых устройств. Эти данные служат основой для обучения ИИ-систем, позволяя им формировать более глубокое представление о внешнем мире и эффективно взаимодействовать с ним [1]. Таким образом, совокупность увеличенных объёмов данных и высоких вычислительных возможностей стала основой для развития более сложных и точных моделей машинного обучения, ускоряя прогресс в области искусственного интеллекта.

Несмотря на некоторые из основных преимуществ, ИИ сталкивается с рядом проблем. Искусственный интеллект обладает потенциалом для преобразования отраслей и повышения благосостояния людей, но он также вызывает опасения по поводу конфиденциальности, этики и потенциальной потери работы. С развитием систем искусственного интеллекта очень важно следить за тем, чтобы при разработке и использовании систем ИИ соблюдались этические нормы.

Предвзятость в системах искусственного интеллекта вызывает серьезную озабоченность, поскольку модели, обученные на основе предвзятых данных, могут привести к постоянной дискриминации и несправедливому обращению [7]. Эти предубеждения заметны в отчетности, отборе и даже групповой атрибуции и скрытой предвзятости, которые впоследствии влияют на пониженные в должности группы и создают неравенство, созданное на протяжении истории.

Системы искусственного интеллекта, применяемые в здравоохранении и финансовом секторе, подвержены уязвимостям, что порождает серьёзные риски для безопасности и конфиденциальности. Такие системы обрабатывают чувствительную информацию — личные данные, финансовые записи и медицинские сведения, — требующую надёжной защиты. Угрозы, способные нанести ущерб моделям ИИ, включают состязательные воздействия,

инверсию моделей и отравление обучающих данных. Особенно уязвимыми являются инфраструктуры здравоохранения, в которых Интернет вещей тесно интегрирован с ИИ-сервисами. В некоторых случаях компрометация этих систем может иметь критические последствия, включая угрозу жизни пациентов. Биомедицинские ИИ-приложения также подвержены атакам на модификацию оборудования и манипуляции привязкой, что ставит под угрозу конфиденциальность и безопасность пациентов.

Развитие технологий в сфере автоматизации, робототехники и ИИ, наблюдаемое в различных отраслях, вызывает обоснованные опасения относительно вытеснения рабочей силы и роста безработицы. Технический прогресс приводит к смещению фокуса занятости, угрожая усилением социального и экономического неравенства. Этап трансформации от концепции индустрии 4.0 к индустрии 5.0 направлен на то, чтобы переосмыслить роль человека и акцентировать внимание на ценности человеческого капитала в условиях высокоавтоматизированной среды [4]. Тем не менее широкое распространение промышленных роботов в производственных цепочках приводит к сокращению рабочих мест и порождает новые вызовы в области охраны труда и соблюдения международных стандартов, касающихся прав человека в сфере занятости, здоровья и безопасности.

Таблица 1. Проблемы и последствия внедрения ИИ в кибербезопасность

ПРОБЛЕМА	ОПИСАНИЕ	ПОСЛЕДСТВИЕ
Доступность данных	Модели ИИ зависят от большого объема данных. В кибербезопасности сбор полных и размеченных наборов данных затруднён из-за чувствительности и конфиденциальности информации.	Неполный и некачественный набор данных может привести к неточной идентификации угроз, увеличению ложных срабатываний и снижению надёжности решений на базе ИИ.
Развивающиеся угрозы	Киберугрозы постоянно меняются, а злоумышленники используют новые методы, которые обходят ранее обученные модели ИИ.	Медленная адаптация ИИ к новым угрозам создаёт уязвимости. Необходима регулярная актуализация данных, что требует ресурсов.
Понимание моделей ИИ	Многие модели ИИ, особенно глубокого обучения, работают как «чёрные ящики» и трудно интерпретируются.	Отсутствие прозрачности решений снижает доверие специалистов к ИИ, особенно в критических ситуациях.
Атаки на ИИ-модели	Злоумышленники могут использовать уязвимости в моделях ИИ, добавляя ложные входные данные.	Это может привести к тому, что система не распознает угрозу и примет неправильное решение, поставив безопасность под угрозу.

Интеграция с существующими системами	Интеграция ИИ в действующую инфраструктуру кибербезопасности может вызывать сложности и проблемы совместимости.	Такое объединение требует изменений в планах развития и может быть затратным. Старые системы могут не поддерживать технологии ИИ.
Ресурсозатратность	Разработка и сопровождение ИИ-систем требует значительных финансовых и вычислительных ресурсов.	Высокая стоимость может ограничить применение ИИ при ограниченном бюджете.
Контроль и соответствие	Нормативная база по ИИ и кибербезопасности всё ещё формируется.	Неопределённость с регулированием мешает внедрению ИИ и может привести к юридическим последствиям.

В таблице 2 приведены основные этические пробелы в кибербезопасности, основанной на ИИ. Предвзятость в моделях ИИ и недостаточное внимание к справедливости создают этические проблемы, в то время как недостаточное внимание к вопросам конфиденциальности может привести к нарушениям прав. Неопределённые рамки подотчетности и непрозрачность систем "черного ящика" подрывают доверие и создают юридические проблемы [3]. Потенциал искусственного интеллекта в области двойного назначения и отсутствие внимания к долгосрочным последствиям для общества, таким как потеря работы или снижение авторитета персонала, усугубляют этические риски и риски безопасности.

Таблица 2. Основные этические пробелы в кибербезопасности, основанной на ИИ

ЭТИЧЕСКИЙ ПРОБЕЛ	ОПИСАНИЕ	ПОСЛЕДСТВИЕ
Предвзятость и справедливость	Обучение моделей ИИ может базироваться на необъективных данных, а сами алгоритмы не всегда учитывают этические последствия.	Это может привести к несправедливому воздействию на отдельные группы и вызвать сомнения в равноправии кибербезопасности.
Конфиденциальность	Часто акцент делается на ИИ как на инструмент защиты конфиденциальности, но игнорируется возможность его злоупотребления в целях слежки.	ИИ может эффективно выявлять угрозы, но при этом нарушать права личности и использоваться для недобросовестного мониторинга.
Ответственность и подотчётность	Чёткое распределение ответственности при ошибках ИИ в кибербезопасности отсутствует.	Возникает правовой и этический пробел, затрудняющий определение виновных при сбоях и инцидентах.

Прозрачность и объяснимость	Использование непрозрачных ИИ-систем в кибербезопасности не сопровождается должным вниманием к этической стороне.	Непрозрачность подрывает доверие и делает трудным этическое оправдание использования ИИ в критических сферах.
Двойное назначение ИИ	Один и тот же ИИ может использоваться как в легитимных, так и в вредоносных целях.	Без должного внимания возрастает риск разработки инструментов, легко доступных злоумышленникам.
Долгосрочные последствия	Большинство исследований сосредоточены на краткосрочных этических проблемах, упуская влияние на общество в будущем.	Игнорирование долгосрочных последствий, таких как утрата рабочих мест и снижение роли человека, ослабляет устойчивость политики.

Ожидается, что в ближайшие годы использование искусственного интеллекта в сфере кибербезопасности будет расширяться благодаря его способности обнаруживать угрозы и реагировать на них.

Эффективное противодействие современным угрозам требует создания принципиально нового формата государственно-частного взаимодействия. Простого обмена информацией об инцидентах и уязвимостях уже недостаточно. Возникает необходимость в построении интегрированных защитных экосистем, в рамках которых ИИ-системы государственных структур и частного сектора действуют согласованно, обмениваясь в реальном времени данными о векторах атак и тактиках реагирования, независимо от отраслевых или юрисдикционных границ.

Особую тревогу вызывает появление качественно нового типа угроз, направленных непосредственно на модели искусственного интеллекта. По мере того как ИИ всё активнее используется для принятия ключевых решений в сфере безопасности, векторы атак смещаются от кражи информации к целенаправленному искажению моделей. Сложные сценарии манипуляции — включая внедрение вредоносных подсказок и искажение обучающих выборок — позволяют злоумышленникам формировать предвзятые или ошибочные поведенческие паттерны внутри ИИ-систем. В результате такие системы продолжают функционировать внешне корректно, но при этом систематически принимают ошибочные решения в критически важных ситуациях [2].

К 2026 году ожидается рост следующих типов атак [10]:

- Состязательные атаки, которые незаметно искажают обучающие данные ИИ.
- Атаки на цепочки поставок, нацеленные на обновление моделей ИИ.
- Эксплойты нулевого дня, специально разработанные для взлома систем безопасности ИИ.

- Атаки с использованием методов социальной инженерии, которые манипулируют алгоритмами обучения ИИ.

- Усовершенствованные методы оперативного внедрения, позволяющие обойти традиционные средства защиты.

К 2026 году ожидается, что программы-вымогатели достигнут нового уровня сложности: киберпреступные группировки начнут активно применять искусственный интеллект и технологии автоматизации для увеличения эффективности и точности атак. Благодаря этим усовершенствованиям вредоносные программы смогут стремительно распространяться по корпоративным сетям, что делает задачу раннего выявления таких угроз особенно критичной. Особую тревогу вызывает рост количества атак, нацеленных на цепочки поставок, поскольку компрометация ключевых поставщиков или партнёров может спровоцировать каскадные сбои в целых отраслях.

В условиях нарастающих рисков организации всё чаще будут обращаться к механизмам киберстрахования как к способу компенсации финансовых потерь от подобных атак, в то время как государственные органы начнут ужесточать регуляторные требования. При этом фишинг по-прежнему остаётся основным каналом доставки программ-вымогателей, а сгенерированные ИИ электронные письма и дипфейковые материалы становятся всё более реалистичными и труднораспознаваемыми.

Одним из наиболее значимых трендов, прогнозируемых на 2026 год, является дальнейшая интеграция ИИ в кибератаки. Уже сейчас искусственный интеллект позволяет злоумышленникам масштабировать свою деятельность, а в ближайшем будущем его влияние только усилится. ИИ-угрозы становятся всё более разнообразными — от персонализированных фишинговых сообщений с грамматически безупречным текстом до продвинутого вредоносного ПО, способного обучаться и обходить средства защиты [5].

По мере того как искусственный интеллект становится всё более распространённым как в личной, так и в профессиональной сфере, растёт обеспокоенность по поводу ненадлежащего использования инструментов ИИ. Одним из самых серьёзных рисков в 2025 году станет утечка данных из-за того, что сотрудники непреднамеренно передают конфиденциальную информацию таким платформам ИИ, как ChatGPT или Google Gemini. Системы ИИ могут обрабатывать огромные объёмы данных, и когда эти данные передаются внешним инструментам ИИ, риск утечки резко возрастает.

Например, сотрудники могут вводить конфиденциальные финансовые данные в инструмент на базе ИИ для создания отчета или проведения анализа, не подозревая, что эти данные могут храниться и потенциально могут быть доступны неавторизованным пользователям.

Интеграция большего числа инструментов на основе искусственного интеллекта в интерфейсы центров мониторинга информационной безопасности (SOC) значительно расширяет возможности специалистов по кибербезопасности. Это позволяет автоматизировать выполнение ключевых задач по обнаружению угроз, существенно снизить уровень ложных срабатываний и повысить оперативность реагирования на инциденты [4].

Социальные сети, охватывающие миллиарды пользователей по всему миру, превратились в приоритетную цель для киберпреступников. К 2026 году ожидается, что сочетание функционала социальных платформ и генеративного искусственного интеллекта будет использоваться для проведения более изощренных и масштабных атак. Используя персонализированные данные и контент, сгенерированный ИИ, злоумышленники смогут реализовывать целевые схемы социальной инженерии, маскировку под доверенные лица и различные виды цифрового мошенничества [3]. Особую опасность представляет не столько каждая из этих технологий сама по себе, сколько их синергия, усиливающая общую степень угрозы.

Преступники всё чаще прибегают к социальным сетям не только для получения конфиденциальной информации, но и в целях манипулирования пользователями, побуждая их к действиям, угрожающим корпоративной безопасности. Особенно уязвимыми в этом контексте являются профессиональные сети, такие как LinkedIn, где под видом легитимного делового контента злоумышленники внедряются в доверительную среду, рассчитывая на ослабление бдительности со стороны работников организаций.

По мере того как всё больше организаций переходят в облако и внедряют устройства Интернета вещей (IoT), поверхность атаки продолжает расширяться. К 2026 году количество устройств Интернета вещей, по прогнозам, превысит 32 миллиарда по всему миру [12]. Безопасность Интернета вещей становится серьёзной проблемой, поскольку злоумышленники могут использовать растущее число подключённых устройств. Многие устройства Интернета вещей, от систем «умного дома» до промышленных датчиков, не имеют надлежащих средств защиты, что делает их привлекательными целями для киберпреступников [6].

Злоумышленники всё активнее осваивают возможности продвинутых ИИ-инструментов для генерации программного кода, переходя от традиционных систем автодополнения, таких как GitHub Copilot, к более мощным платформам, способным создавать полноценное вредоносное ПО по единичному запросу. Такие технологии значительно

упрощают процесс разработки вредоносных программ, снижая порог вхождения в киберпреступную деятельность.

В ближайшие годы ситуация с дефицитом квалифицированных специалистов по информационной безопасности будет оказывать серьёзное влияние на способность организаций противостоять возрастающим угрозам. Несмотря на постоянный рост инвестиций в сферу кибербезопасности и расширение линейки защитных решений, нехватка компетентных кадров для настройки, обслуживания и интеграции этих систем приведёт к разрозненности и снижению общей эффективности средств защиты [11]. При этом чрезмерная зависимость от множества внешних поставщиков без необходимого внутреннего экспертного ресурса лишь усугубит уязвимость: управлять такой сложной экосистемой станет затруднительно, а её устойчивость — заметно ниже.

Заключение

В 2025 году ситуация в сфере кибербезопасности определяется ростом числа атак с использованием искусственного интеллекта и растущей уязвимостью платформ социальных сетей, чтобы справиться с этими проблемами, организациям необходимо инвестировать в защиту на основе искусственного интеллекта и внедрить подход «Никому не доверяй» для обеспечения безопасности в облаке и Интернете вещей. Кроме того, компаниям необходимо подготовиться к ужесточению нормативно-правовой базы и растущей потребности в киберстраховании. Киберпреступность развивается беспрецедентными темпами, и компании, которые не смогут адаптироваться, рискуют стать следующей жертвой. Сейчас самое время действовать, чтобы защитить цифровые активы и обеспечить безопасность в будущем.

Список использованных источников:

1. Aadil, G. M. & Maaz, S.M. (2024). Human-AI-Collaboration in Healthcare. International Journal for Research in Applied Science & Engineering Technology (IJRASET).
2. Banjanovic-Mehmedovic, L. & Husakovi'c, A. (2023). Edge AI: Reshaping the Future of Edge Computing with Artificial Intelligence. Basic technologies and models for implementation of Industry 4.0 Conference. <https://doi.org/10.5644/PI2023.209.07>
3. Borenstein, J., & Howard, A. (2020). Emerging Challenges in AI and the Need for AI Ethics Education, AI and Ethics. 1 (2021). 61 - 65. <https://doi.org/10.1007/s43681-020-00002-7>
4. Camacho, N. G. (2024). The Role of AI in Cybersecurity: Addressing Threats in the Digital Age. Journal of Artificial Intelligence General Science (JAIGS) ISSN:3006-4023, 3(1), 143–154.

5. Chen, Y. (2020). IoT, Cloud, Big Data and AI in Interdisciplinary Domains. <https://doi.org/10.1016/j.simpat.2020.102070>.
6. Dave, D. M., & Mandvikar, S. (2023). Augmented Intelligence: Human-AI Collaboration in the Era of Digital Transformation. *International Journal of Engineering Applied Sciences and Technology* 8, 24–33. <https://doi.org/10.33564/IJEAST.2023.v08i06.003>
7. Jawaid, S. A. (2023). Artificial Intelligence with Respect to Cyber Security. *Journal of Advances in Artificial Intelligence*. 1(2), 96–102. <https://doi.org/10.18178/JAAI.2023.1.2.96-102>
8. Kent, M. D. & Kopacek, P. (2021). Do We Need Synchronization of the Human and Robotics to Make Industry 5.0 a Success Story?. *International Symposium for Production Research*. 302 – 311.
9. Li, F. (2024). Application and Challenges of Artificial Intelligence in Cybersecurity. *Applied and Computational Engineering*. 47 (1), 262–268. <https://doi.org/10.54254/2755-2721/47/20241480>.
10. Mahadik, S., Ravat, N. J., Singh, K. Y. & Yadav, S. D. (2020). Feature Analysis Detection Algorithms-Review Paper. *International Journal of Advanced Research in Science, Communication and Technology*, 194–199. <https://doi.org/10.48175/IJARSC-645>
11. Pan, K. (2024). Ethics in the Age of AI: Research of the Intersection of Technology, Morality, and Society. *Lecture Notes in Education Psychology and Public Media*. 40 (1). 259–262 <https://doi.org/10.54254/2753-7048/40/20240816>
12. Rizvi, M. (2023). Enhancing Cybersecurity: The Power of Artificial Intelligence in Threat Detection and Prevention. *International Journal of Advanced Engineering Research and Science*. 10 (5), 055–060. <https://doi.org/10.22161/ijaers.105.8>