

УДК 33:004.8**DOI: 10.34670/AR.2026.13.65.087****Актуальные проблемы применения искусственного интеллекта в хозяйственной деятельности экономических субъектов****Соколовский Александр Александрович**

Кандидат экономических наук,
доцент кафедры экономической безопасности,
Российская академия народного хозяйства
и государственной службы при Президенте Российской Федерации,
119571, Российская Федерация, Москва, пр-кт Вернадского, 82;
e-mail: sokolovskiy-aa@ranepa.ru

Грибов Павел Геннадьевич

Кандидат экономических наук,
доцент кафедры экономической безопасности,
Российская академия народного хозяйства
и государственной службы при Президенте Российской Федерации,
119571, Российская Федерация, Москва, пр-кт Вернадского, 82;
e-mail: gribov-pg@ranepa.ru

Аннотация

Данная статья посвящена изучению актуальных вопросов регулирования деятельности хозяйствующих субъектов в условиях развития технологий искусственного интеллекта. Отмечается, что отсутствие выработанного как глобального, так и регионального подхода в области регулирования подобных технологий создает значительные сложности по выработке механизма обеспечения экономической безопасности хозяйствующих субъектов всех уровней, а также построения эффективной системы управления.

Для цитирования в научных исследованиях

Соколовский А.А., Грибов П.Г. Актуальные проблемы применения искусственного интеллекта в хозяйственной деятельности экономических субъектов // Экономика: вчера, сегодня, завтра. 2026. Том 16. № 4А. С. 902-909. DOI: 10.34670/AR.2026.13.65.087

Ключевые слова

Искусственный интеллект, экономическая безопасность, цифровые технологии, робототехника, система управления, правовое регулирование, кибербезопасность.

Введение

Повсеместная цифровизация социально-экономических процессов и использование искусственного интеллекта (ИИ) открывают значительные перспективы для повышения эффективности национальной экономики за счёт автоматизации процессов, оптимизации ресурсов, создания новых бизнес-моделей и стимулирования инноваций. Эти возможности проявляются в различных сферах — от промышленности и здравоохранения до государственного управления и транспорта, но особенно следует выделить преимущества, которые несет с собой цифровизация и использование ИИ в социальной сфере.

Однако стремясь максимизировать возможные результаты, которые дает повсеместное использование цифровых технологий и искусственного интеллекта, не следует забывать о том, что неконтролируемое применение ИИ способно привести к формированию и развитию ряда существенных угроз, затрагивающих различные сферы жизни — от кибербезопасности и экономики до корпоративных рисков и социальных последствий. Злоумышленники могут использовать цифровые технологии и ИИ для проведения продвинутых фишинговых атак, генерации различных вредоносных кодов, обхода существующих систем защиты информационных систем предприятий. Кроме того, не следует преуменьшать вероятность реально существующих рисков «отравления» обучающих систем: злоумышленники могут намеренно внедрить искажённую информацию в наборы данных, используемые для обучения моделей искусственного интеллекта, что приведёт к предвзятым или ошибочным решениям.

Основная часть

В современной экономике, искусственный интеллект зачастую используется для создания реалистичного рода дипфейков — поддельных аудио-, видео- и фотоматериалов, которые чрезвычайно сложно отличить от настоящих. Это усиливает угрозы для персональной безопасности, способствует распространению дезинформации и может влиять как на различные социальные, так и на происходящие в различных странах мира политические процессы [Елифанова, 2024]. Негативные последствия непродуманного повсеместного внедрения цифровых технологий и ИИ в финансовой сфере вполне реально может стать причиной манипуляций на фондовых рынках, спекулятивного поведения инвесторов, распространения фальшивых аналитических заключений и финансовых рекомендаций и др. Необдуманное использование тех или иных неотработанных, непроверенных алгоритмов реально способно воспроизводить и усиливать искажения в данных, что ведёт к дискриминации при кредитовании и страховании, ограничению доступа тех или иных экономических агентов к финансовым услугам [Воробьева, 2022]. Еще одной угрозой непродуманного внедрения цифровых технологий и искусственного интеллекта, является существующая зависимость от ограниченного круга поставщиков языковых моделей, которая способна привести к формированию и развитию существенных рисков для стабильности финансовых систем, поскольку сбой или атака на таких поставщиков может дестабилизировать финансовый рынок.

Использование ИИ на промышленных предприятиях становится безусловным трендом. Следует отметить, что сегодня сформированы отдельные области, в которых применение подобных технологий стало логичным развитием накопленного опыта использования цифровых технологий предыдущего поколения. К таким сферам можно отнести предиктивное

обслуживание оборудования, которое начиная с 2000-х гг. активно развивалось на базе систем удаленного мониторинга и контроля. Экономическая эффективность подобных систем была обоснована на основе сокращения сроков простоя оборудования и прогнозирования выхода из рабочего состояния наиболее ценных единиц оборудования промышленных предприятий. Автоматизированные системы контроля качества на промышленных предприятиях стали активно использоваться еще в 1990-х гг., однако, начиная с 2010-х гг., подобные системы стали активно использоваться в паре с системами компьютерного зрения. Безусловно, наиболее ярко применение цифровых технологий и их развитие на базе ИИ получили при создании цифровых двойников. Подобная технология, позволила повысить эффективность работы масштабных промышленных предприятий, а также при анализе состояния работы инфраструктурных объектов.

Оптимизация логистики в торговой сфере, а также автоматизированное построение маршрутов транспортных предприятий стало одним из первых применений цифровых систем. Представить работу современных сервисов такси и служб доставок без использования цифровой картографии и систем распределения заказов уже невозможно.

Повсеместно внедряя цифровые технологии и искусственный интеллект, не следует забывать о том, что ИИ-системы часто требуют доступа к большим объемам данных, включая личную или конфиденциальную информацию. Все это существенно повышает риск утечек конфиденциальных данных, несанкционированного доступа к информации, составляющей коммерческую тайну. К примеру, сотрудники, загружая в чат-боты корпоративные данные (бизнес-презентации, стратегические планы, те или иные финансовые документы компании и т. д.), непреднамеренно создают скрытые каналы утечки конфиденциальной информации. Внедряя цифровые технологии и искусственный интеллект, следует иметь в виду и то, что существует также риск того, что искусственный интеллект непреднамеренно запоминая и воспроизводя части обучающих данных, включая персональные или конфиденциальные сведения, может в соответствии с запросом, поступившем от злоумышленников, выдать им информацию, составляющую коммерческую тайну.

Наибольшую опасность для реализации технологий ИИ и киберфизических систем, представляет отсутствие единых стандартов и механизмов контроля использования цифровых технологий и искусственного интеллекта на международном уровне. Отсутствие единых международных стандартов и механизмов контроля в сфере использования цифровых технологий и ИИ, создаёт условия для его бесконтрольного и возможно злонамеренного использования по нескольким ключевым направлениям. Это связано с фрагментарностью регулирования, сложностями в определении ответственности, технологическим разрывом, геополитическими факторами и другими существенными аспектами. Характеризуя такую уязвимость, как фрагментарность регулирования, следует отметить, то обстоятельство, что разные страны реализуют на практике различные подходы к регулированию ИИ — от строгого законодательного контроля до практически полного отсутствия какого бы то ни было регулирования.

В Европейском Союзе с 2024 г. действует Закон об искусственном интеллекте (Artificial Intelligence Act), направленный на защиту людей от рисков, связанных с ИИ. В США регулирование использования ИИ, происходит на уровне штатов и отдельных ведомств, при этом акцент делается на поддержку инноваций, а не на жёсткий контроль за процессом

использования ИИ. В некоторых странах Азии, как например в Китае, где использования ИИ характеризуется сочетанием стимулирования инноваций с жёстким государственным контролем, особенно в сферах, связанных с национальной безопасностью, этикой и общественным порядком. С 2023 года в Китае были утверждены «Временные меры по управлению генеративными системами искусственного интеллекта» [Регулирование искусственного интеллекта: первые шаги, www].

В США на федеральном уровне до настоящего времени не существует единого федерального закона об использовании искусственного интеллекта (ИИ). Регулирование строится на добровольных стандартах и инициативе штатов. Национальный институт стандартов и технологий (NIST) выпустил фреймворк управления рисками ИИ (AI RMF), но он носит рекомендательный характер. С 2023 штаты начали разрабатывать собственные законы, а компаниям рекомендуется использовать «NIST AI RMF» для отчётности. При этом в некоторых штатах (например, Колорадо) приняты локальные законы, регулирующие высокорискованные системы ИИ, но общая картина остаётся разрозненной.

Япония, до недавнего времени практически не имела комплексного законодательства в сфере искусственного интеллекта (ИИ) и, лишь в 2025 года приняла «Закон о содействии исследованиям, разработкам и использованию технологий, связанных с ИИ», который начал действовать в 2025 году. Этот закон ориентирован на поддержку инноваций, а не на жёсткое регулирование рисков использования ИИ. В нём отсутствуют штрафные санкции, классификация систем по уровням риска и обязательные механизмы правоприменения. Подход Японии основан на стимулировании сотрудничества между государством, бизнесом и академическими кругами, а не на жёстком контроле [Богомазов, 2024].

Великобритания до недавнего времени опиралась на существующие правовые рамки (законы о защите данных, ответственности за качество продукции), а не на отдельное законодательство об искусственном интеллекте (ИИ). И лишь в 2023 году правительство Великобритании, опубликовало «Белую книгу» с планами будущего регулирования, но акцент делался на гибкости и добровольном соблюдении пяти принципов: безопасность, прозрачность, справедливость, подотчётность и состязательность. Обязательства по этим принципам на первых порах не были обязательными.

Наблюдаемое ужесточение применения ИИ-систем в развитых экономиках, компенсируется наличием стран, где контроль за использованием искусственного интеллекта минимален. Это позволяет экономическим агентам переносить деятельность в юрисдикции с менее строгим регулированием использования ИИ, минимизируя ответственность и контроль. Разработчики могут избегать соблюдения стандартов безопасности, этики и защиты данных, выбирая страны с лояльным законодательством.

Отсутствие единых стандартов и механизмов контроля использования ИИ-систем на международном уровне несёт угрозу невозможности установления ответственных за те или иные ошибки и злоупотребления, связанные с ИИ. Следует отметить, что искусственный интеллект часто воспринимается как «чёрный ящик», в связи с чем, достаточно сложно установить, кто несёт ответственность за ошибки или злоупотребления. В отсутствие международных норм и чётких правовых рамок ответственность за действия искусственного интеллекта может оставаться неопределённой, что затрудняет привлечение к ответственности разработчиков, пользователей или владельцев систем. Существующая практика

свидетельствует, что, если искусственный интеллект, разработанный в одной стране, наносит ущерб в другой (например, через формирование и распространение дезинформации или организацию и проведение кибератак), определить юрисдикцию для иска, доказать причинно-следственную связь и взыскать компенсацию чрезвычайно сложно. Кроме того, международный механизм принудительного исполнения подобных решений отсутствует.

Отсутствие общих критериев риска и единой международной нормативно-правовой базы использования цифровых технологий и ИИ, отсутствие единого списка «запрещённых» применений искусственного интеллекта (например, социальный скоринг, массовое распознавание лиц), различие подходов к допустимой степени автономности систем, требованиям к тестированию безопасности, правилам информирования пользователей о взаимодействии с ИИ, ведут к тому, что страны и организации по-своему интерпретируют риски, что создает определенные возможности для злонамеренного использования ИИ.

Различные международные организации (ООН, ЮНЕСКО, ОЭСР) вырабатывают и распространяют определенные документы, связанные с практическими аспектами использования искусственного интеллекта, однако эти документы, носят рекомендательный характер и не имеют законодательной силы. В современных условиях, единый глобальный надзорный орган за использованием ИИ, по сути, отсутствует, как отсутствует и система мониторинга нарушений, и механизмы санкций за несоблюдение добровольных стандартов использования искусственного интеллекта. Все это позволяет странам и компаниям игнорировать общепринятые этические принципы использования ИИ и минимизировать результаты любых международных усилий по регулированию использования ИИ.

Сохраняющаяся анонимность интернета и трансграничный характер распространения данных позволяют распространять вредоносные ИИ-инструменты (фишинговые чат-боты, генераторы поддельных документов), проводить атаки с использованием ИИ без идентификации злоумышленников, обходить национальные блокировки через VPN и децентрализованные сети, а отсутствие международных механизмов контроля за киберпространством, усиливает эти риски [Мочалов, 2021].

Практические последствия всего вышеперечисленного имеют в основном негативный характер. Прежде всего следует отметить, что отсутствие международного регулирования регулированию использования ИИ способствует:

- распространению дезинформации и дипфейков, поскольку генеративные модели способны создавать правдоподобные изображения, аудио- и видеоконтент, который, как правило, используется для манипуляций, мошенничества и других подобных целей;
- кибератакам и утечкам конфиденциальных данных, поскольку незащищённые ИИ-системы становятся мишенями для хакеров, а отсутствие стандартов безопасности облегчает кражу данных;
- избыточному регулированию применения систем мониторинга и контроля на основе ИИ-систем не только в экономической деятельности, но и в социальной сфере;
- геополитической напряжённости, так как различия в подходах к регулированию использования искусственного интеллекта между странами усиливают взаимное недоверие конкуренцию и реально способно послужить причиной формирования и развития международных конфликтов.

Минимизация рисков использования искусственного интеллекта требует разработки и

заключения международных договоров с обязательными требованиями к прозрачности ИИ-систем, оценке рисков до практического внедрения ИИ создания разрешительно-запретительных механизмов позволяющих минимизировать злонамеренное использование искусственного интеллекта (например, выдача и отзыв лицензий) и др. Поэтому, развитие международного сотрудничества в сфере использования искусственного интеллекта крайне важно. Это сотрудничество, можно развивать в рамках ОЭСР, ООН и Глобального партнёрства по искусственному интеллекту (GPAI), что способствовало бы выработке общих стандартов и этических принципов. Именно в рамках международных организаций, возможна разработка и реализация на практике, международных норм и стандартов регулирования практического использования ИИ, усиление киберзащиты, контроль за использованием данных, программы переквалификации кадров и этический контроль за применением ИИ-технологий.

Повсеместное непродуманное внедрение искусственного интеллекта, чрезмерное доверие к ИИ-системам без критического осмысления их решений, способно привести к критическим ошибкам и непредвиденным последствиям. Существует опасность, что люди могут «потерять себя», всецело доверившись виртуальному разуму, что повлияет на их самостоятельность и способность к критическому мышлению.

Насущная необходимость практической минимизации критических рисков и угроз, вызываемых к жизни повсеместным внедрением технологий искусственного интеллекта в социально-экономическую практику, заставляет вспомнить, сформулированные еще более половины столетия назад, писателем Айзеком Азимовым, вымышленные законы (принципы) этического кодекса робототехники. Эти принципы, названные Азимовым «законами робототехники», впервые были сформулированы писателем в рассказе «Хоровод» в 1942 году и никогда не имели юридической силы, как и применения на практике. Однако они оказали существенное влияние на представления о необходимости формирования этики машин и дискуссии о необходимости обеспечения безопасности людей от искусственного интеллекта. Таких законов (принципов) Айзек Азимов сформулировал три:

1. Робот не может причинить вред человеку или своим бездействием допустить, чтобы человеку был причинён вред;
2. Робот должен подчиняться всем приказам человека, если они не противоречат первому закону;
3. Робот должен защищать своё существование в той мере, в которой это не противоречит первому и второму законам [Три закона робототехники Айзека Азимова, www].

Иерархия законов (принципов) построена так, что ценность человеческой жизни ставится выше всего остального. Робот не может оправдать свои действия ссылкой на приказ или на собственное выживание, если это причиняет вред человеку.

Переход вопросов контроля и управления робототехникой в экономическую плоскость произошел лишь с развитием систем искусственного интеллекта (ИИ). В 2020 г. Ф. Паскуале были сформулированы дополнительные законы робототехники, со следующим содержанием:

- ИИ должен дополнять профессионалов, а не заменять их.
- ИИ не должен подделывать человека.
- В области ИИ следует предотвращать усиление гонки вооружений.
- Роботы и ИИ должны быть приписаны к своим создателям или владельцам, чтобы всегда можно было определить ответственного [Соколов, 2023].

Заключение

Развитие понимания функционирования искусственного интеллекта приведет к встраиванию подобных систем в экономический механизм функционирования предприятий. Очевидно, что экономическая ценность результата работы ИИ-систем может быть измерена в денежном обращении и вовлечена в экономический оборот, что требует не только ее нормативно-правового регулирования, но и создания инфраструктуры для проведения коммерческих операций.

Библиография

1. Богомазов Н.Л. Тенденции развития искусственного интеллекта в Японии: «Общество 5.0» и статус «Электронное лицо» // Гуманитарные ведомости ТГПУ им. Л.Н. Толстого. 2024. №2. С.56-64.
2. Воробьева Е.А., Блажевич О.Г. Цифровые технологии в финансовой сфере: целесообразность их применения в трансформации экономики Российской Федерации // Научный Вестник: Финансы, банки, инвестиции. 2022. №1. С. 18-25.
3. Епифанова Т.В., Копейкин К.И. Проблемы законодательного регулирования объектов, созданных с использованием дипфейк-технологий в России и за рубежом // Северо-Кавказский юридический вестник. 2024. №3. С. 121-127.
4. Мочалов А. Н. Парадокс анонимности в Интернете и проблемы ее правового регулирования // Вестник Сургутского государственного университета. 2021. № 4. С. 111-121.
5. Регулирование искусственного интеллекта: первые шаги. URL: <https://foresight.hse.ru/mirror/pubs/share/886272540.pdf>
6. Соколов Д. В. К синтезу человеческой мудрости и вычислительной мощности. Рецензия на книгу Ф. Паскуале «Новые законы робототехники. Апология человеческих знаний в эпоху искусственного интеллекта» // Управление наукой: теория и практика. 2023. Т. 5, № 1. С. 238-243.
7. Три закона робототехники Айзека Азимова. URL: <https://enjoy-robotics.ru/tri-zakona-robototekhniki>

Current Problems of Applying Artificial Intelligence in the Economic Activity of Business Entities

Aleksandr A. Sokolovskii

PhD in Economics,
Associate Professor, Department of Economic Security,
Russian Presidential Academy
of National Economy and Public Administration,
119571, 82, Vernadskogo ave., Moscow, Russian Federation;
e-mail: sokolovskiy-aa@ranepa.ru

Pavel G. Gribov

PhD in Economics,
Associate Professor, Department of Economic Security,
Russian Presidential Academy of
National Economy and Public Administration,
119571, 82, Vernadskogo ave., Moscow, Russian Federation;
e-mail: gribov-pg@ranepa.ru

Sokolovskii A.A., Gribov P.G.

Abstract

This article is devoted to the study of current issues of regulating the activities of business entities in the context of the development of artificial intelligence technologies. It is noted that the lack of a developed both global and regional approach in the field of regulating such technologies creates significant difficulties in developing a mechanism for ensuring the economic security of business entities at all levels, as well as building an effective management system.

For citation

Sokolovskii A.A., Gribov P.G. (2026) Aktual'nye problemy primeneniya iskusstvennogo intellekta v khozyaystvennoy deyatel'nosti ekonomicheskikh sub'ektov [Current Problems of Applying Artificial Intelligence in the Economic Activity of Business Entities]. *Ekonomika: vchera, segodnya, zavtra* [Economics: Yesterday, Today and Tomorrow], 16 (4A), pp. 902-909. DOI: 10.34670/AR.2026.13.65.087

Keywords

Artificial intelligence, economic security, digital technologies, robotics, management system, legal regulation, cybersecurity.

References

1. Bogomazov, N.L. (2024). Tendentsii razvitiya iskusstvennogo intellekta v Yaponii: «Obshchestvo 5.0» i status «Elektronnoye litso» [Trends in the Development of Artificial Intelligence in Japan: "Society 5.0" and the Status of "Electronic Person"]. *Gumanitarnyye vedomosti TGPU im. L.N. Tolstogo*, (2), 56-64.
2. Epifanova, T.V., & Kopeikin, K.I. (2024). Problemy zakonodatel'nogo regulirovaniya obyektov, sozdannykh s ispolzovaniyem dipfeyk-tekhnologiy v Rossii i za rubezhom [Problems of Legislative Regulation of Objects Created Using Deepfake Technologies in Russia and Abroad]. *Severo-Kavkazskiy yuridicheskiy vestnik*, (3), 121-127.
3. Mochalov, A.N. (2021). Paradoks anonimnosti v Internetе i problemy yeye pravovogo regulirovaniya [The Paradox of Anonymity on the Internet and Problems of Its Legal Regulation]. *Vestnik Surgutskogo gosudarstvennogo universiteta*, (4), 111-121.
4. Regulirovaniye iskusstvennogo intellekta: pervyye shagi [Regulation of Artificial Intelligence: First Steps]. (n.d.). Retrieved from <https://foresight.hse.ru/mirror/pubs/share/886272540.pdf>
5. Sokolov, D.V. (2023). K sintezu chelovecheskoy mudrosti i vychislitel'noy moshchnosti. Retsenziya na knigu F. Paskuale «Novyye zakony robototekhniki. Apologiya chelovecheskikh znaniy v epokhu iskusstvennogo intellekta» [Toward a Synthesis of Human Wisdom and Computational Power. Review of F. Pasquale's Book "New Laws of Robotics. Apology of Human Knowledge in the Age of Artificial Intelligence"]. *Upravleniye naukoy: teoriya i praktika*, 5(1), 238-243.
6. Tri zakona robototekhniki Ayzeka Azimova [Isaac Asimov's Three Laws of Robotics]. (n.d.). Retrieved from <https://enjoy-robotics.ru/tri-zakona-robototekhniki>
7. Vorobyeva, E.A., & Blazhevich, O.G. (2022). Tsifrovyye tekhnologii v finansovoy sfere: tselesoobraznost ikh primeneniya v transformatsii ekonomiki Rossiyskoy Federatsii [Digital Technologies in the Financial Sphere: The Feasibility of Their Application in the Transformation of the Economy of the Russian Federation]. *Nauchnyy Vestnik: Finansy, banki, investitsii*, (1), 18-25.