

<https://doi.org/10.22363/2313-2337-2025-29-2-524-543>
EDN: ZNDTRO

Научная статья / Research Article

Возможности правового нивелирования политических и организационно-управленческих рисков применения искусственного интеллекта

А.С. Киселев  

Федеральное государственное образовательное бюджетное учреждение высшего образования «Финансовый университет при Правительстве Российской Федерации»,
г. Москва, Российская Федерация
alskiselev@fa.ru

Аннотация. Рассматриваются основные риски применения технологий, основанных на элементах искусственного интеллекта, а также правовые меры, которые рационально было бы предпринять для минимизации этих рисков. Установлено, что использование технологий, основанных на отдельных элементах искусственного интеллекта (далее – ИИ), в политической сфере несет с собой как потенциал для улучшения эффективности управленческой работы, так и серьезные риски. Подчеркивается отсутствие у искусственного интеллекта человеческой морали, и вытекающие отсюда этические проблемы применения ИИ, проблема сохранения конфиденциальности данных, с которыми работает ИИ, проблема создания и распространения ложной информации (на примере дипфейков), проблема дискриминации искусственным интеллектом и неравного доступа граждан к технологиям ИИ и др. Вполне очевидно, что прямо сейчас не создан сильный ИИ, имеющий в своем арсенале набор возможностей, который максимально приблизит его к человеку по когнитивным способностям. Тем не менее, тенденции последних лет говорят о том, что общество, государство, даже сами разработчики не всегда могут предвидеть негативные последствия новых технологий. В этой связи представляется наиболее целесообразным выстроить механизм опережающего регулирования, который будет учитывать самые ключевые риски внедрения технологий, основанных на ИИ. В качестве правовых мер минимизации указанных рисков предложено создать и контролировать соблюдение новых правовых норм, создать как государственные, так и независимые экспертные органы и комиссии по контролю за ИИ и этике искусственного интеллекта, обучать граждан, и особенно чиновников, работающих с ИИ, не только необходимым навыкам, но так же и вопросам этики использования ИИ, особое внимание уделить принципам прозрачности и открытости в применении этих новых технологий. Указывается на необходимость учитывать зарубежный опыт и создать международные этические и юридические нормы, призванные снизить риски, сопутствующие распространению технологий искусственного интеллекта.

Ключевые слова: государство, искусственный интеллект, право, риски применения ИИ, политические риски, управленческие риски, правовое регулирование ИИ, опережающее регулирование, общество

Конфликт интересов. Автор заявляет об отсутствии конфликта интересов.

© Киселев А.С., 2025



This work is licensed under a Creative Commons Attribution 4.0 International License
<https://creativecommons.org/licenses/by-nc/4.0/legalcode>

Поступила в редакцию: 22 августа 2024 г.

Принята к печати: 15 апреля 2025 г.

Для цитирования:

Киселев А.С. Возможности правового нивелирования политических // *RUDN Journal of Law*. 2025. Т. 29. № 2. С. 524–543. <https://doi.org/10.22363/2313-2337-2025-29-2-524-543>

The Possibilities of Legal Leveling of Political, Organizational and Managerial Risks of Using Artificial Intelligence

Alexander S. Kiselev  

Federal State Educational Budgetary Institution of Higher Education "Financial University under the Government of the Russian Federation", *Moscow, Russian Federation*

alskiselev@fa.ru

Abstract. The article examines the primary risks associated with the application of artificial intelligence (AI) technologies and legal measures to minimize these risks. It establishes that the use of AI technologies in the political sphere presents both potential for improved efficiency in governance and significant risks. The absence of human morality in artificial intelligence is emphasized, leading to ethical concerns regarding AI application, data privacy issues, the creation and spread of misinformation (for example, deepfakes), AI-driven discrimination, and unequal access to AI technologies. While a “strong AI” with cognitive abilities closely matching humans does not currently exist, recent trends suggest that society, governments, and even developers may not always foresee the negative consequences of emerging technologies. Therefore, proactive regulation is crucial to address key risks associated with the implementation of AI technologies. Proposed legal measures to minimize these risks include creating and enforcing new legal norms, establishing governmental and independent expert bodies to oversee AI and AI ethics, educating citizens (especially officials working with AI) on necessary skills and ethical considerations, and prioritizing transparency and openness in the application of these technologies. The article also emphasizes the need to consider international experience and develop international ethical and legal norms to reduce the risks associated with the proliferation AI technologies.

Key words: state, artificial intelligence, law, risks of AI application, political risks, governance risks, AI legal regulation, proactive regulation, society

Conflict of interest. The author declares no conflict of interest.

Received: 22th August 2024

Accepted: 15th April 2025

For citation:

Kiselev, A.S. (2025) The Possibilities of Legal Leveling of Political, Organizational and Managerial Risks of Using Artificial Intelligence. *RUDN Journal of Law*. 29 (2), 524–543. (in Russian). <https://doi.org/10.22363/2313-2337-2025-29-2-524-543>

Введение

Искусственный интеллект (далее – ИИ), бурно развивающийся в последние годы, имеет широкий спектр возможностей и перспектив в сфере управления. Уже в середине XX века ЭВМ использовались для автоматизации повторяющихся

монотонных задач, для обработки данных, составления отчетов, аналитики и т.д. В этом состояла, например, суть проекта ОГАС (Общегосударственная автоматизированная система учета и обработки информации), который реализовывался в СССР согласно идеям академика В.М. Глушкова.

Сегодня похожие проекты, где-то в меньших, где-то в больших масштабах, но на новом техническом уровне, реализованы во многих странах мира. Управление с помощью информационно-телекоммуникационных технологий позволяет руководителям разных уровней и сотрудникам фокусироваться на более сложных и стратегических важных задачах. Этот процесс начался в XX веке, а сегодня автоматизированное управление становится необходимым, как воздух, для любого государства. Во многих экономических и производственных процессах просто невозможно участвовать, если в стране не произошла интеграция ИКТ-технологий в медицине, науке, искусстве, образовании, промышленности, а, прежде всего, в политике и управлении государством (Novitsky, 2020:48-49).

Современные исследователи посвящают свои научные работы специфике трансформации политики посредством внедрения технологий, основанных на элементах ИИ. Так, С.В. Володенков изучает возможности перевода политики в цифровой формат, цифровую инфраструктуру гражданско-политического активизма. Аналогичные вопросы поднимаются в научных исследованиях ученых О.А. Башевой, С.Е. Козлова, К. Кравфорд, Р. Кало. Исследование Е.М. Ильиной посвящено политике и управлению в условиях цифровой трансформации. Другая группа исследователей, среди которых В.П. Казарян, К.Ю. Шутова, Е.Д. Сорокина, Р. Юсти, С. Геринг, поднимают проблему существования этических рисков. Инвестиционным рискам при внедрении ИИ уделено большое внимание в работе Н.А. Новицкого. Риски противоправного применения ИИ отражены в исследованиях С.В. Лемайкиной, А.М. Постарнак, Н.Н. Черногора. О.А. Ястребов и Ю.А. Гаврилова в своих научных трудах указывают, что регулировать ИИ необходимо для обеспечения безопасности социума. Можно сделать заключение о том, что в доктрине широко обсуждаются вопросы как перспектив применения ИИ, так и рисков, которые сопутствуют этому процессу.

Методологическая основа. Используются общенаучные методы, такие как индукция и дедукция, анализ и синтез, абстрагирование, обобщение, системный метод, а также частно-научные методы – формально-юридический и сравнительно-правовой.

Цель исследования – формирование предложений по правовому нивелированию политических и организационно-управленческих рисков применения искусственного интеллекта.

Задачами, служащими раскрытию цели, являются: выявление основных рисков применения ИИ на сегодняшний день в политике и в управлении; оценка их потенциальной опасности; изучение зарубежного опыта организационно-правового преодоления существующих и потенциальных рисков; формулировка предложений по минимизации и нивелированию выявленных рисков.

Искусственный интеллект в политике: вчера и сегодня

Невозможно представить современную внутреннюю и внешнюю политику без прогнозирования грядущих событий и детальной аналитики уже произошедших.

Современные программы, основанные на работе элементов искусственного интеллекта, имеют возможность быстро анализировать большие объемы данных и обнаруживать скрытые паттерны и тенденции. Это, несомненно, помогает в принятии более конструктивных управленческих решений и прогнозировании дальнейших действий. К примеру, для выявления закономерностей и прогнозирования результатов президентских или парламентских выборов. В частности, «пермскими учеными за полтора года до президентских выборов 2008 г. в России была спрогнозирована победа Д.А. Медведева, когда его личность как политика была еще не была широко известна» (Пуина, 2022:407–408).

Искусственный интеллект и машинное обучение расширяют потенциал технологий моделирования «цифровых двойников (от англ. digital twin) – виртуальных аналогов любых физических объектов или процессов (например, ЭБ-карта Сингапура или трехмерная модель левых камер сердца пациента от Philips Heart Model). Эти технологии становятся ядром системы «интеллектуальных двойников» (Intelligent Twins) – новой открытой архитектуры для интеллектуальной трансформации государственных органов и городских служб (интеллектуальный двойник города), отраслей промышленности (интеллектуальный двойник промышленности) и предприятий (интеллектуальный двойник бизнеса), способной воспринимать трехмерные объекты, взаимодействовать в любых областях, выносить точные решения, непрерывно эволюционировать, совместно с другими программами и людьми работать в любых сферах деятельности, выносить точные суждения и непрерывно эволюционировать, обеспечивая интеллектуальный опыт для людей, городов будущего и предприятий во всех сценариях»¹.

При этом, стоит отметить, что под интеллектуальными двойниками понимается инновационная система поэтапной модернизации государственных органов власти и частных компаний, предприятий, в основе которой лежат технологии, основанные на элементах искусственного интеллекта. Иными словами, посредством ИИ создаются механизмы, способные в режиме реального времени воспринимать все изменения в законодательстве, экономике, социальной сфере, выносить точные суждения по разным вопросам, работать с любыми типами информации, а что самое главное – эволюционировать и подстраиваться под запросы государства и общества.

Перспективы применения технологий, основанных на элементах искусственного интеллекта, в политике

ИИ может предоставлять рекомендации и советы на основе анализа данных и изучения предыдущих решений во многих сферах жизнедеятельности. С.В. Володенков указывает, что «Цифровые решения, интерфейсы, алгоритмы и механики большинством специалистов рассматривались преимущественно как надстройки по отношению к традиционным общественно-политическим процессам. Однако в определенный момент стало очевидно, что цифровизация не является лишь надстройкой, расширением, дополнением традиционного общественно-политического пространства. Интенсивное развитие цифровых технологий и их активное внедрение в ключевые сферы жизнедеятельности государства и общества привело к формированию

¹ Intelligent Twins. Совместное создание интеллектуальных двойников и построение мира интеллектуальных технологий. Huawei, IDC, 2020. Р. 28. Режим доступа: <https://www.huawei.ru/intelligent-twins/> (дата обращения: 01.08.2024).

принципиально новых моделей, принципов и механизмов цифрового общественно-политического взаимодействия, появлению новых акторов, способных конкурировать с традиционными политическими игроками и эффективно влиять на массовое сознание» (Volodenkov, 2022: 341).

Стоит отметить возможности автоматизированного анализа рисков с помощью ИИ. Специализированные программы, функционируя в штатном режиме, позволяют обнаруживать и предсказывать потенциальные риски и угрозы для крупного, среднего и малого бизнеса, помогая как в управлении отдельным предприятием, так и целыми отраслями экономики. При больших вычислительных мощностях и скоростях обработки данных возможно создать системы, позволяющие оценивать риски в режиме реального времени при управлении всем государством, соответственно, хозяйственная роль суперкомпьютеров и программ, основанных на ИИ, сильно возрастет в ближайшее десятилетие. Анализ рисков позволяет как отдельным компаниям, так и чиновникам принимать соответствующие меры для снижения рисков и повышения эффективности управления. Для планирования этих мер опять же, можно использовать возможности ИИ.

О подобных изменениях говорит Г. Ловинг, отмечая, что «эффекты, которыми сопровождается распространение современных цифровых технологических, и следующая из этого трансформация традиционных общественно-политических институтов и процессов, самым непосредственным образом влияют на содержательные, структурные и функциональные параметры жизнедеятельности государства и общества, а также их ключевых институтов, оказывают непосредственное влияние на существующий политический порядок, систему традиционных властных отношений и общественно-политическую динамику в целом» (Loving, 2019).

Еще одной позитивной тенденцией является автоматизация процесса обслуживания клиентов в сфере государственных и коммерческих услуг. Ранее данные новшества начали появляться в работе колл-центров коммерческих предприятий и банков. Затем МФЦ и отдельные органы власти также начали заимствовать положительный опыт, передавая искусственному интеллекту некоторые несложные стандартизированные функции сотрудников. Например, чат-боты и виртуальные помощники могут предоставлять быстрые и точные ответы на стандартные вопросы обращающихся граждан. Стоит отметить, что порой в процессе оказания подобных услуг возникают трудности, однако алгоритмы доказали свою эффективность. Их применение в процессе оказания консультационных услуг развивается, так как предыдущий опыт оценивается как положительный.

В работе органов власти программы, основанные на элементах ИИ, могут помочь в планировании и управлении ресурсами, в оптимизации рабочих процессов, таких как снабжение и логистика. Это позволяет как коммерческим организациям, так и органам власти сократить издержки, увеличить скорость работы.

Перспективы использования искусственного интеллекта в управлении очень обширны. С развитием технологий ИИ ожидается, что они будут играть все более значимую роль в повышении производительности труда, эффективности и инновационности в управленческой практике. Уже сейчас зарубежные исследователи отмечают, что ИИ может получить доступ к автономному регулированию сферы энергетики и иных общественно важных систем (Ren Hu, 2022: 447–448). В этой связи, как отмечают исследователи, «особый интерес представляет сегодня исследование цифровых корпораций как новых посредников между гражданами и властью, способных

влиять на уровень включенности первых в процессы обсуждения и принятия политических решений. Кроме того, цифровые корпорации, как показывает анализ актуальной практики, могут сегодня выступать в качестве мощного фактора, влияющего на возможности и форматы демократического участия граждан в значимых социально-политических процессах» (Volodenkov, 2021:101).

Рассматривая результаты применения новых технологий, стоит помнить, что технологии, как правило, нужны, чтобы облегчить труд человека, а не усложнить его. Поэтому, при внедрении любой технологии, стоит учитывать, как с ней будет взаимодействовать пользователь, обращающийся за получением услуги, чиновник, занимающийся вопросами их оказания населению, работник, компании, директор и многие другие участники процесса цифрового взаимодействия. Необходима всесторонняя оценка и учет интересов лиц, которым предстоит столкнуться на практике с новыми технологиями в сфере управления. Это значит, что не стоит сразу вводить новшество, которое, по ключевым показателям, приносит прибыль и оценивается как высокоэффективный инструмент. На первом месте всегда должен стоять человек, его ментальное и физическое здоровье, его права и законные интересы. В противном случае, без учета данных факторов, человек может стать «заложником» цифровизации, невольным участником процессов, которые претят нормам гуманизма и человеческой природе. Поэтому, для таких технологий требуется длительный процесс тестирования и обучения, запуск пилотных проектов, работа в «тестовых» режимах, оценка общественностью и специальными рабочими группами и т.п. Данные меры позволяют снизить риски внедрения искусственного интеллекта сегодня и в будущем. Одновременно с этим следует также учитывать этические и социально-бытовые аспекты использования ИИ, чтобы обеспечить безопасность его применения. Но все мы, в любом случае, должны привыкнуть к тому, что ИИ уже сейчас является неотъемлемой частью нашей жизни.

Тема политического использования искусственного интеллекта вызывает широкое обсуждение среди исследователей в сферах политики, экономики и права. Несомненно, ИИ предлагает огромный потенциал в улучшении эффективности политических процессов и в принятии решений, однако с его использованием связаны и некоторые очевидные и неочевидные риски.

Риски применения технологий, основанных на искусственном интеллекте

Первый риск, который следует учитывать при внедрении ИИ в системы политического управления, – это потеря человеческого контроля. ИИ, будучи основанным на алгоритмах и машинном обучении, может лишь имитировать понимание норм морали, этики и эмоций, что на первый взгляд может казаться преимуществом. Однако политика включает в себя множество нюансов, которые практически невозможно учесть с помощью алгоритмов. От человека, занимающийся политикой, часто требуется эмоциональный подход к решению проблем населения. Особенно опасно отсутствие человеческой морали, например, при применении вооружений, оснащенных ИИ. Обыкновенное оружие контролируется людьми: политиками и военными, и без их воли не применяется. Автономные «умные» системы вооружений представляют глобальную опасность именно в силу отсутствия у ИИ человеческой морали.

Такие системы уже стоят на вооружении ряда наиболее развитых стран, и уже принимали решения атаковать людей (Kazaryan & Shutova, 2023: 350–351). Регулярно высказываются мнения, что такое вооружение должно быть запрещено, как особо опасное. Например, такой позиции придерживается действующий генеральный секретарь ООН².

Ученые еще в конце XX в. доказали, что без развитого эмоционального интеллекта сложно строить коммуникации как в сфере управления коммерческими организациями, так и в сфере политики. Эмоциональный интеллект имеется только у человека, он основан на понимании своих и чужих эмоций, контроле и управлении ими. Человек может, с помощью обращения к эмоциям, доказать состоятельность той позиции, которую зачастую нельзя донести только логикой. Если политический процесс будет полностью зависеть от ИИ, то потеря человеческого фактора может привести к отсутствию эмпатии, сочувствия в процессе управления, которые необходимы для создания связи между политиками и населением страны. Таким образом, ИИ может предлагать идеи, но их воплощение и процесс взаимодействия с людьми, учет их мнений, должен осуществлять исключительно человек.

Второй риск связан с проблемой безопасности и конфиденциальности данных. Использование ИИ неразрывно связано со сбором и анализом больших объемов данных, в том числе, персональных. В политической сфере, где точная информация является важнейшим инструментом при принятии решений, возникает риск злоупотребления данными. Используя возможности ИИ, злоумышленники могут создавать средства и инструменты для взлома паролей, распознавания уязвимостей и обхода систем безопасности для получения доступа к конфиденциальной информации, данным тысяч и миллионов пользователей, зарегистрированных на государственных и муниципальных порталах.

Нарушение конфиденциальности может привести к неправомерному использованию информации, манипуляциям и подрыву демократических ценностей. С помощью алгоритмов машинного обучения и анализа больших данных ИИ может быть использован различными политическими силами для манипулирования мнением общественности.

Е.С. Козлов подтверждает данную мысль, указывая, что «распространение социально-политических видео в интернет-среде действительно может приводить к реальным действиям и прогрессу в отношении стратегических целей различных кампаний» (Kozlov, 2020:160). Или же, например, ИИ может использоваться для создания качественно сделанных фальшивых новостей на основе дипфейков (реалистичной замены изображения и звука, в том числе лиц и голоса, с помощью ИИ) или манипулирования информацией в социальных сетях, чтобы влиять на общественное мнение и результаты выборов (Kiselev, 2021:56–57; Lemaykina, 2023:144). Так, например, в 2018 году в США распространялось видео-дипфейк, на котором предыдущий президент негативно высказывался о действующем. Это была просто демонстрация возможностей подобных технологий. Однако стоит ожидать, что в ближайшие годы этот новый вид «черного пиара» будет активнейшим образом использоваться в политической борьбе, и в мошеннических коммерческих схемах (Kiselev, 2021:58).

² Ranger S. Killer AI robots must be outlawed, says UN chief // Официальный сайт «ZDnet». 2018. Режим доступа: <https://www.zdnet.com/article/killer-ai-robots-must-be-outlawed-says-un-chief/> (дата обращения: 04.11.2024).

Стоит отметить, что с помощью ИИ также можно создавать автоматизированные системы для распространения дезинформации и проведения информационных атак на конкретных политиков или на всю политическую систему целиком. Например, ИИ может использоваться для создания ботов, которые распространяют фейковую информацию и поддерживают определенные политические мнения, создавая иллюзию поддержки или протестов. В 2019 году в США распространялось дипфейковое видео с якобы пьяным спикером Палаты представителей (Lemauskina, 2023:144). Уже есть множество примеров из разных стран, когда с помощью дипфейков мошенники успешно обманывали банки и граждан с целью доступа к их финансам.

Еще в 2021 г. с помощью ИИ злоумышленники изменили голос и обманули представителей банка в ОАЭ. Позвонивший «директор» потребовал сделать перевод на сумму в 35 млн долл. После звонка на электронную почту банка пришло письмо от компании и от юриста, что убедило сотрудников банка в легитимности операции³. После инцидента меры предосторожности стали повышать, тем не менее технологии тоже не стоят на месте, как и лица, которые их применяют в своих противоправных схемах. Например, в 2024 г. в Гонконге сотруднику крупной финансовой компании якобы позвонил по видеоконференцсвязи финансовый директор из Лондона и дал задание перечислить на указанные счета деньги на общую сумму более 25 млн долл.⁴ Дело в том, что образ руководителя фирмы был сделан посредством технологии Deepfake настолько реалистично, что работник не заметил подвоха и выполнил операцию, обман раскрылся лишь спустя некоторое время⁵. Известным элементом современной политики являются розыгрыши пранкеров, которые звонят известным политикам и другим гражданам и, имитируя чью-либо речь, подводят собеседника к неосторожным высказываниям. Инструменты по созданию дипфейков открывают бесконечный простор для подобных спекуляций. В результате в ближайшем будущем подлинность любой информации будет вызывать всеобщие сомнения.

В будущем, с ускорением обработки информации и совершенствованием таких инструментов, как чат GPT и его аналогов, возможна автоматизация процессов дезинформации в круглосуточном режиме. Программы, основанные на самообучении, могут следить, каким новостям доверяет подавляющее большинство пользователей социальных сетей и зрителей телеканалов и генерировать на основе собранных данных новости, включая как текст, так и любое изображение в таком контексте, чтобы максимальное количество пользователей поверили в это.

Сегодня (в 2024 г.) тексты, составляемые ИИ, не всегда похожи на те, которые написал бы реальный человек, но нейросети очень быстро учатся и уже могут писать рассказы, письма, длинные тексты, и даже стихи, с учетом стиля письменной речи конкретных авторов (писать официальные письма в стиле Петра Первого, стихи в стиле Пушкина, Лермонтова, рассказы в стиле Зощенко, Беляева, а также современных деятелей искусства, политики, героев фильмов и т.д.). Чат-боты GPT и его

³ Bank Robbers in the Middle East Reportedly 'Cloned' Someone's Voice to Assist with \$35 Million Heist. Официальный сайт «Gizmodo». 2021. Режим доступа: <https://gizmodo.com/bank-robbers-in-the-middle-east-reportedly-cloned-someone-1847863805> (дата обращения: 04.11.2024).

⁴ Прим.: ок. 200 миллионов гонконгских долларов.

⁵ A company lost \$25 million after an employee was tricked by deepfakes of his coworkers on a video call: Police. Официальный сайт «Business Insider». 2024. Режим доступа: <https://www.businessinsider.com/deepfake-coworkers-video-call-company-loses-millions-employee-ai-2024-2> (дата обращения: 04.11.2024).

аналоги доступны бесплатно любому пользователю Интернета через поисковые системы и в соцсетях: дело двух минут открыть подобную страницу и попросить ИИ написать текст на заданную тему в стиле конкретного писателя, поэта, персонажа. Социальные сети все больше заполняются вполне качественным по форме контентом, созданным искусственным интеллектом. Возможно, недалек уже тот день, когда никто из нас не сможет определить, писал ли конкретный пост ИИ или человек, и только специальные программы смогут вычислять это, обладая всеми возможностями ИИ, недоступными человеку. Именно поэтому стоит сосредоточить внимание не на препятствовании развитию ИИ (что в принципе уже невозможно на сегодняшний день), а технологиям, которые могут «разоблачать» ИИ, находя закономерности, свойственные только алгоритмизированной машине.

Рассмотрим данный риск с другой стороны. Используя ИИ-системы анализа поведения, правительства стран по всему миру могут отслеживать и изучать в режиме реального времени действия граждан в социальных сетях, в том числе обходя режимы анонимности. О такой проблеме, в частности, писал С. Джадали, отмечая, что действия миллионов пользователей отслеживаются благодаря установленным расширениям в интернет-браузерах Chrome и Firefox. Специалист приходит к следующему выводу: «Если в момент установки расширение не передает и не анализирует информацию об интернет-серфинге пользователя, нет никаких гарантий, что этого не произойдет в будущем»⁶. Примечательно, что в 2020 г. Апелляционный суд девятого округа США признал факт массового сбора и анализа метаданных, включая информацию о телефонных переговорах граждан, что нарушает Конституцию США⁷. Решение было вынесено спустя семь долгих лет после того, как Эдвард Сноуден обнародовал прямые доказательства прослушивания телефонных переговоров спецслужбами США. Помимо этого, каждому человеку свойственен определенный стиль устной и письменной речи. Поэтому даже скрывающего свои персональные данные и местонахождение пользователя можно с высокой долей вероятности идентифицировать по стилю речи. Чем больше будет образцов текста, которые сможет обработать ИИ, тем точнее будет результат анализа.

Очевидно, что в связи с вышеописанным, под угрозой находятся, закрепленные во многих правовых системах, права на свободу информации, тайну частной корреспонденции и личной жизни, авторское право и т.д. Эти вопросы станут актуальны в ближайшем будущем, так как пока ИИ не настолько совершенен, чтобы следить постоянно за всем и каждым. Но решать их нужно уже сегодня. Следует учитывать различные варианты того, как ИИ может быть использован для наблюдения за общественной активностью, незаконным сбором персональных данных и проведением массовой слежки за населением.

Третий риск – это потенциальные этические проблемы. ИИ может столкнуться с такими дилеммами в процессе принятия решений, которые мы относим к области морали. До сих пор нет четкого ответа на следующие вопросы: какими критериями руководствуется и должен руководствоваться ИИ в выборе наиболее эффективных решений? Или: как программа определит достоинства и приоритеты политических

⁶ Jadali S. DataSpии: The catastrophic data leak via browser extensions. Сайт “SecurityWithSam”. Режим доступа: <https://securitywithsam.com/2019/07/dataspii-leak-via-browser-extensions/> (дата обращения: 05.11.2024).

⁷ United States Court of Appeals for The Ninth Circuit, No. 13-50572, 2020. Режим доступа: <https://cdn.arstechnica.net/wp-content/uploads/2020/09/nsa-ruling.pdf> (дата обращения: 05.11.2024).

целей? Вполне логично, что ИИ должен руководствоваться базовыми правилами социальной справедливости, которые характерны для конкретной страны и общества, где ИИ может значительно повлиять на распределение ресурсов или принятие решений, основываясь на информации, которая не вызовет сомнений и нареканий у большинства граждан. Соответственно, здесь стоит проблема создания этического кодекса искусственного интеллекта, о которой ведутся дебаты как в нашей стране, так и за рубежом. В апреле 2021 г. «Европейская комиссия представила проект Регламента о гармонизации правил, применяемых к системам ИИ (далее – Регламент), в основу которого был положен риск-ориентированный подход; в октябре 2021 г. в России крупнейшие технологические компании, университеты, исследовательские центры и фонды подписали Кодекс этики в сфере ИИ⁸ (далее – Кодекс); в ноябре 2021 г. был принят важный международный документ – Рекомендации ЮНЕСКО в сфере этики ИИ⁹» (Sorokova, 2022: 159).

Центральное место в Рекомендациях ЮНЕСКО занимают четыре основных принципа, закладывающие основу для систем ИИ: уважение прав и свобод человека, мирное сосуществование, обеспечение инклюзивности и благополучие окружающей среды, экосистем мира. Схожие принципы заложены в отечественном Кодексе: человеко-ориентированный и гуманистический подход, уважение автономии и свободы воли человека, соответствие закону, недискриминация, оценка рисков и т.д.

Вполне очевидно, что наш Кодекс соответствует большей части положений, закрепленных в Рекомендациях ЮНЕСКО в сфере этики ИИ. В США, например, нет закона или этического кодекса, который бы вводил единые требования ко всем системам, основанным на ИИ. Во времена президентства Б. Обамы в 2016 г. была создана дорожная карта развития ИИ¹⁰, в ней, прежде всего, внимание уделено перспективам технологического развития государства, и лишь малая часть уделялась вопросам этики ИИ. Тем не менее, такие компании как Amazon, Google, Apple, IBM и Microsoft в скором времени создали некоммерческое партнерство, одна из ключевых целей которого состояла в разработке и продвижении общих стандартов этики ИИ для коммерческих компаний – участников соглашения¹¹. Развитие вопросов этики ИИ на сегодняшний день в США происходит фрагментарно.

Примечателен опыт Китая: подавляющее внимание сосредоточено на обеспечении информационной безопасности как части национальной безопасности, в частности, защите персональных данных, немалый акцент сделан на недопущении

⁸ Кодекс этики в сфере ИИ (создан Альянсом в сфере искусственного интеллекта, Аналитическим центром при Правительстве РФ, Минэкономразвития РФ и др.) // Официальный сайт «ethics.a-ai». 2021. Режим доступа: <https://ethics.a-ai.ru/> (дата обращения: 04.11.2024).

⁹ Recommendation on the Ethics of Artificial Intelligence (Рекомендации об этике искусственного интеллекта), UNESCO. Национальный портал в сфере искусственного интеллекта (ИИ) и применения нейросетей в России «ai.gov.ru». 2021. Режим доступа: https://ai.gov.ru/knowledgebase/dokumenty-mezhdunarodnykh-organizatsiy-po-ii/2021_rekomendaciya_ob_etike_iskusstvennogo_intellekta_recommendation_on_the_ethics_of_artificial_intelligence_unesco/?ysclid=m332zmgms77785714662 (дата обращения: 04.11.2024).

¹⁰ The Obama Administration's Roadmap for AI Policy. Сайт Harvard Business Review. 2016. Режим доступа: <https://hbr.org/2016/12/the-obama-administrations-roadmap-for-ai-policy> (дата обращения: 05.11.2024).

¹¹ Amazon, Google, Facebook, IBM, and Microsoft form AI non-profit. Сайт «ZDnet». 2016. Режим доступа: https://translated.turbopages.org/proxy_u/en-ru.ru.7c676982-672a0683-3533a4a8-74722d776562/https/www.zdnet.com/article/amazon-google-facebook-ibm-and-microsoft-form-ai-non-profit/ (дата обращения: 05.11.2024).

недобросовестной конкуренции посредством ИИ. Выделим базовые принципы этики ИИ в Китае: люди должны иметь право применять ИИ или отказаться от его применения, ИИ должен быть направлен на обеспечение благосостояния граждан, ИИ не должен использоваться в незаконной деятельности, ИИ должен обеспечивать справедливость, невмешательство в частную жизнь граждан Китая¹². В целом можно сказать, что принципы этики ИИ в Китае являются самыми подробными и конкретными на сегодняшний день.

С 2019 г. в России и на западе было принято несколько документов, призванных ускорить развитие ИИ и обеспечить снижение возможных при этом рисков. О необходимости этического подхода к внедрению ИИ говорили как российские, так и западные политики. Хотя европейские и российские документы различаются не только содержанием, но и характером применения, они имеют схожие черты: человекоцентрированность, гуманистичность, различение степеней риска систем ИИ, опора на экспертные научные знания. Значительное внимание в них уделяется защите персональных данных и выстраиванию системы взаимодействия акторов ИИ. Разница заключается, в первую очередь, в том, что в ЕС существуют уровни общеевропейского законодательства и множества различных национальных законодательств, которые необходимо согласовывать¹³. В России же законодательство едино на территории всей страны. Этические нормы применения ИИ в Евросоюзе, как замечает Е.Д. Сорокова, демонстрируют эволюцию от «мягких» к более обязывающим (Sorokova, 2022:160). Вероятно, эта тенденция будет только усиливаться (как следствие увеличения рисков применяемого ИИ), и будет наблюдаться и в других правовых системах.

Отметим, что без хотя бы минимального «этического кодекса ИИ», внедрение его в управление людьми настолько опасно, что становится вообще нецелесообразным.

Четвертым риском, который мы считаем необходимым выделить, является усиление социального неравенства в обществе. Исследования, проводившиеся в США, показали, что алгоритмы Google дискриминируют женщин, предлагая им менее оплачиваемые вакансии по сравнению с мужчинами (что может быть верным отражением перекоса на рынке труда, а не в алгоритмах Google), а также темнокожих – распознавая в их фотографиях обезьян (горилл) и предлагая им более криминальную информацию (Yuste & Goering, 2017:159–161). Были и более резонансные случаи. Например, Американский союз защиты гражданских прав (прим.: правозащитная некоммерческая организация) в 2019 г. подал жалобу в полицию Детройта в связи с тем, что система распознавания лиц, функционирующая на основе ИИ, определила добропорядочного и не имеющего судимости афроамериканца Роберта Уильямса как разыскиваемого преступника. Юристы союза предоставили доказательства о том, что система много лет предвзято относилась к людям определенной расы, делая существенные ошибки. Даже в полицейском участке отметили, что лицо настоящего преступника совершенно не похоже на лицо задержанного. Через две недели после

¹² Волков К. В Китае выпустили кодекс этических принципов для искусственного интеллекта // Официальный сайт «Российская газета». 2021. Режим доступа: <https://rg.ru/2021/10/04/v-kitae-vypustili-kodeks-eticheskikh-principov-dlia-iskusstvennogo-intellekta.html> (дата обращения: 05.11.2024).

¹³ A European approach to artificial intelligence. Официальный сайт Европейской Комиссии. 2021. Режим доступа: <https://digital-strategy.ec.europa.eu/en/policies/european-approach-artificial-intelligence> (дата обращения: 01.08.2024).

ареста Уильямса освободи¹⁴. Однако стоит учесть, что гражданину был нанесен существенный репутационный ущерб.

Так же отмечалась неэффективность ИИ-систем, использующихся в некоторых штатах в судебных процессах (Crawford & Calo, 2016:311). Если доступ к используемому ИИ будет ограничен и доступен только богатым корпорациям, государству и отдельным гражданам, это может привести к укреплению их власти и усилению неравенства, конфликтности в обществе. Помимо этого компания, имеющая возможность использовать передовые технологии, основанные на ИИ, и конкретные политики могут использовать ИИ в своих интересах, в том числе коррупционных, поскольку разработка систем контроля за самыми передовыми технологиями всегда отстает от развития и применения этих технологий. Это верно и для искусственного интеллекта. Если к основанным на элементах ИИ технологиям, которые значительным образом могут повлиять на нашу жизнь, не будут иметь доступ все слои общества, то это может усилить уже существующие социальные, управленческие и иные проблемы, продуцируя раскол между социальными группами и слоями населения. Так, например, большая часть статей про искусственный интеллект в крупнейшем западном научном журнале «Nature» – это статьи про использование ИИ в медицине. Технологии искусственного интеллекта позволяют более точно распознавать болезни на ранних стадиях, создавать новые методы лечения и лекарственные препараты. Естественно, что при существующей экономической модели все новейшие достижения будут доступны только нескольким наиболее богатым процентам населения. Но ведь и среди тех, кто может лишен доступа к новейшим достижениям в силу нехватки денежных средств, есть и очень полезные для общества и государства люди. По нашему мнению, доступ к лучшим социальным благам должен быть у каждого человека, т.к. здоровое и образованное население может дать лучшие результаты во всех видах деятельности, чем бедное, больное и неграмотное.

Кроме того, рост неравенства в обществе может привести к росту недовольства, что может ухудшить жизнь всем. В российской истории есть прекрасный пример М.В. Ломоносова – человека из низшего сословия, который, получив доступ к науке и поддержку государства, принес немалую пользу стране. Если открыть перед каждым путь Ломоносова, то, вероятно, и «Ломоносовых» станет больше. Более того, развитие технологий ИИ открывает перспективы, еще недавно казавшиеся совершенно фантастическими: не только излечения многих сложных генетических и психических болезней, но совершенно новые способы обмена информацией и управления физическими структурами силой мысли и т.п. (Yuste & Goering, 2017:159–163). С другой стороны, если новейшие технологии останутся привилегией ограниченного круга людей, то у них может возникнуть желание использовать свой доступ к этим инструментам противоправно.

Следующий риск связан с технической стороной внедрения ИИ, а именно с недостатком прозрачности и объяснимости принятия решений искусственным интеллектом. Этот риск актуален в текущих условиях, поскольку в некоторых судах США начали активно применять ИИ в правосудии. Одна из таких систем искусственного интеллекта получила название «Система оценки общественной безопасности» (Public Safety Assessment, PSA). Она была разработана частным Фондом Лауры и

¹⁴ ACLU of Michigan Complaint re Use of Facial Recognition. Сайт ACLU. 2020. Available at: <https://www.aclu.org/documents/aclu-michigan-complaint-re-use-facial-recognition> (дата обращения: 05.11.2024).

Джона Арнольдс. Среди факторов, по которым анализируется личность: возраст, характер предъявленных обвинений, наличие криминальной биографии и т.д. (всего 9). Кроме того, PSA присваивает каждому проверяемому рейтинг общественной опасности – от 1 до 6. Оценка предназначена для того, чтобы помочь судьям оценить уровень риска для обвиняемого в случае освобождения из-под стражи в ожидании суда по уголовному обвинению»¹⁵. Судьи штата Нью-Джерси, Техаса, Юта и некоторых других штатов уже начали использовать PSA в своей работе. Причем для широких кругов общественности не раскрываются принципы работы данной системы. Главная опасность заключается в том, что алгоритмы могут «отпустить на волю» опасного преступника, а судья доверится системе и примет волевое решение. Последствия, вне всякого сомнения, могут быть очень трагичные.

Многие алгоритмы ИИ работают на основе сложных математических моделей, трудных для понимания среднестатистическим человеком. Часто даже разработчики конкретного ИИ не могут объяснить, как и почему ИИ принимает именно те решения, которые он принимает. Следовательно, возникает опасность неправильного или изначально предвзятого принятия решений без возможности объяснения причины таких действий. Это может стать проблемой, особенно когда требуется дать пояснения о принятых решениях общественности или заинтересованным гражданам. В то же время возможны случаи сбоев и систематических ошибок, которые не поддаются прогнозированию.

Итак, можно констатировать, что использование искусственного интеллекта и технологий, основанных на его отдельных элементах, в политической сфере несет с собой как потенциал для улучшения эффективности управленческой работы, так и серьезные риски, такие как утрата человеческого фактора в управлении, проблемы безопасности и конфиденциальности данных, этические вопросы и усиление социального неравенства, недостаток прозрачности и объяснимости решений, принимаемых искусственным интеллектом. Поэтому внедрение ИИ должно осуществляться с учетом этих и многих других рисков, при активном участии представителей широкой общественности, в целях обеспечения баланса между автоматизацией процессов, ответственностью государства и общества и паритетом интересов всех и каждого.

Все эти примеры подчеркивают важность этического и ответственного использования ИИ в политике и управлении. Именно поэтому применение искусственного интеллекта именно в политике (где часто практикуется обман, коррупция, манипуляции общественным мнением, засекречивание информации, тайные договоренности, и т.д.) и в управлении может быть сопряжено с многочисленными рисками, такими как технические несовершенства ИИ, возможности потенциального нарушения прав человека и гражданина.

Организационно-правовые меры и предложения по минимизации и нивелированию рисков от внедрения ИИ в сферу управления государством

Как отмечают В.П. Казарян и К.Ю. Шутова «Законы робототехники Айзека Азимова более полувека оставались фантастикой, но сегодня юридическое регулирование ИИ и его использование в самой юриспруденции становится повседневной

¹⁵ Utah Public Safety Assessment Frequently Asked Questions. Официальный сайт The Journal Branch of Utah. Режим доступа: <https://www.utcourts.gov/en/court-records-publications/publications/court-publications/court-reports/psa/faq.html> (дата обращения: 05.11.2024).

практикой» (Kazaryan & Shutova, 2023: 353). Именно поэтому отсутствие закрепленных правил этики и мер ответственности за незаконное применение ИИ, включая злоупотребление ИИ для достижения политических и управленческих целей, значительно повышает воплощение этих рисков в реальность. Если речь идет о менее опасных и незначительных рисках применения ИИ, то возможно регулирование посредством издания подзаконных нормативно-правовых и локальных нормативных актов. К примеру, когда речь идет о внедрении ИИ в процессы принятия решений на уровне муниципальных органов власти или отдельного предприятия. С помощью правовых инструментов можно значительно минимизировать эти риски и обеспечить более безопасное и эффективное использование ИИ в государственном управлении и в политике.

Считаем, что для минимизации рисков необходимо разработать и применять соответствующее законодательство. Нормативная база должна обеспечивать ответственность за неправомерное использование ИИ, нарушения базовых конституционных прав человека и гражданина, исключая любое вредоносное вмешательство в личную жизнь граждан. В законодательстве также следует предусмотреть механизмы контроля за системами ИИ, чтобы гарантировать их безопасность и соответствие этическим нормам. Действующее законодательство в области регулирования ИИ¹⁶ предусматривает принципы работы с технологиями ИИ, цели и задачи внедрения таких технологий, обязанности субъектов и многое другое. В законах и подзаконных нормативно-правовых актах отмечается особая опасность систем, работающих на основе технологий ИИ, однако до сих пор не принят ни один закон, предусматривающий ответственность за неправомерное использование ИИ. Считаем, что необходимо законодательно предусмотреть меры ответственности за самые опасные для национальной безопасности противоправные деяния.

В Российской Федерации с 2024 г. действуют отдельные нормы, предусматривающие административную ответственность за применение дипфейков в определенных случаях. Укажем, что в п. 20 Постановления Пленума Верховного Суда РФ от 25 июня 2024 г. № 17 «Привлечению к административной ответственности по части 1 статьи 5.12 КоАП РФ подлежат также заказчики и лица, выполнившие работы, оказавшие услуги по созданию (подготовке) агитационных материалов всех видов с использованием изображений любого физического лица, в том числе умершего (фотографий, видеозаписей или произведений изобразительного искусства, изображений, созданных с использованием компьютерных технологий)...»¹⁷. Более того,

¹⁶ Федеральный закон от 31 июля 2020 г. № 258-ФЗ «Об экспериментальных правовых режимах в сфере цифровых инноваций в Российской Федерации» (ред. от 8 августа 2024 г. № 223-ФЗ) // Российская газета. 2020. № 173; Федеральный закон от 24 апреля 2020 г. № 123-ФЗ «О проведении эксперимента по установлению специального регулирования в целях создания необходимых условий для разработки и внедрения технологий искусственного интеллекта в субъекте Российской Федерации – городе федерального значения Москве и внесении изменений в статьи 6 и 10 Федерального закона «О персональных данных» (ред. от 8 августа 2024 г. № 233-ФЗ) // Российская газета. 2020. № 92; Указ Президента РФ от 10 октября 2019 г. № 490 «О развитии искусственного интеллекта в Российской Федерации» (ред. от 15 февраля 2024 г. № 124) // Собрание законодательства Российской Федерации. 2019. №41. Ст. 5700; Указ Президента РФ от 9 мая 2017 г. № 203 «О Стратегии развития информационного общества в Российской Федерации на 2017–2030 годы» // Собрание законодательства Российской Федерации. 2017. № 20. Ст. 2901.

¹⁷ Постановление Пленума Верховного Суда Российской Федерации от 25 июня 2024 г. № 17 «Об отдельных вопросах, возникающих у судов при рассмотрении дел об административных правонарушениях, посягающих на установленный порядок информационного обеспечения выборов и референдумов» // Бюллетень Верховного Суда Российской Федерации. 2024. № 8.

мошенники начинают использовать дипфейки для создания голосов умерших родственников для совершения преступлений¹⁸. К сожалению, не предусматривается ответственность за оскорбление посредством создания дипфейка, хотя потенциальный ущерб может быть очень высоким, гораздо выше, чем «обычное» оскорбление. На наш взгляд, следует предусмотреть административную и уголовную ответственность за создание и распространение дипфейков, основных на образах мировых лидеров, знаменитостей и, разумеется, любого человека.

Отсюда следует, что в российском законодательстве тема дипфейков еще не проработана: отсутствуют не только конкретные нормы, но даже и необходимые определения. Все больше исследователей, среди которых следует отметить Т.Я. Хабриеву, Н.Н. Черногора, которые говорят о необходимости регулирования искусственного интеллекта в различных сферах человеческой жизнедеятельности (Khabrieva, 2023:5–12, Chernogor, 2023:5–15). Л.В. Бертовский делает вывод о формировании «высокотехнологичного права, внутренней стороной которого является использование высоких технологий для решения задач, появляющихся в процессе правоприменения, а внешней – регулирование общественных отношений, возникающих при их использовании» (Bertovsky, 2021:735).

На наш взгляд, требуется создание как государственных, так и независимых экспертных органов и комиссий по этике искусственного интеллекта. Это поможет в оценке и контроле использования ИИ в сфере политики и управления как на общегосударственном уровне, так и на местах. Эти органы должны формироваться из высококвалифицированных специалистов, которые смогут анализировать и оценивать все виды рисков, включая вопросы технологического применения, и рекомендовать меры по их полному нивелированию или минимизации.

Кроме того, считаем целесообразным проведение обучения государственных и муниципальных служащих всех органов власти с применяемыми в их работе новыми технологиями. При этом в процессе прохождения ими образовательных программ должны быть раскрыты вопросы этики применения искусственного интеллекта, возможные риски его использования и иные проблемы, которые мы отразили в настоящем исследовании. Государственные органы власти в лице Правительства РФ и Министерства высшего образования и науки РФ на федеральном уровне должны, на наш взгляд, поддерживать различные образовательные программы, в том числе курсы и семинары, способствующие осведомлению сотрудников органов власти о правовых и этических аспектах использования ИИ как в текущей деятельности, так и в недалеком будущем.

Для расширения доступа к информации о применении ИИ в политике принципы прозрачности и открытости играют важнейшую роль. Открытые научные дискуссии, публикации данных и отчетов о федеральных и региональных проектах, связанных с ИИ, помогут снизить риски неприятия гражданами технологий, основанных на ИИ, а также риски принятия недемократичных решений. Представители гражданского общества, к примеру, независимые эксперты и группы общественного контроля за внедрением отдельных технологий в сфере управления, основанных на ИИ, должны иметь возможность участвовать в обсуждении и в непосредственном контроле применения ИИ. По оценкам О.А. Башаевой, «развитие современных социальных платформ послужило толчком к формированию сетевого сообщества и

¹⁸ Аферисты начали просить у россиян деньги под видом умерших родственников // Сайт Life. Режим доступа: <https://life.ru/p/1700173?ysclid=m34d2h07u3374638819> (дата обращения: 05.11.2024).

появлению нового типа гражданских онлайн-приложений, основанных на самостоятельном решении гражданским обществом целого ряда проблем (отслеживание качества воздуха в городах или образования стихийных свалок, ликвидация последствий стихийных бедствий и прочее» (Bashaeva, 2020:46). Аналогичные платформы можно создавать в целях осуществления общественного контроля за ИИ.

В целях минимизации рисков использования искусственного интеллекта в сфере политики и управления мы предлагаем закрепление в законах и иных нормативно-правовых актах следующих норм:

- о прозрачности использования и применения искусственного интеллекта (алгоритмы и системы искусственного интеллекта должны быть понятными и объяснимыми, чтобы их действия мог проверить любой человек, жизнь которого затрагивают такие технологии, включая решения, которые влияют на политические и управленческие процессы в государстве на любом уровне; с этой целью с 2018 г. в США создается объяснимый искусственный интеллект (Explainable Artificial Intelligence – XIA), тем не менее, пока ни в одной стране мира нет закона, регламентирующего его применение);

- о запрете дискриминации и обеспечении социального и политического равенства (никто не должен быть ограничен в использовании высоких технологий, основанных на ИИ – вне зависимости от пола, расы, социального положения, равно как и подвергаться предвзятой оценке со стороны технологий ИИ). Как в ситуации с Робертом Уильямсом в США, любые системы поиска преступников, основанные на обработке биометрии ИИ, должны проверяться человеком, в противном случае может пострадать репутация гражданина, а доверие к правоохранительным органам будет подорвано. Примечательно, что пока только в этических рекомендациях и кодексах прописывают принцип недискриминации ИИ, но в конкретных нормах (УК РФ, КОАП РФ) данный принцип еще не нашел отражения;

- о неприкосновенности человеческих прав и свобод (Это означает, что искусственный интеллект не должен принимать решения, которые нарушают основные права и свободы, включая запрет на обработку и анализ персональных данных, разрешения на обработку которых человек не давал). К примеру, в ст. 7 Федерального закона от 24 апреля 2020 г. № 123-ФЗ указано, что системы ИИ должны собирать только обезличенные данные. Однако требуется по аналогии с законом Евросоюза об ИИ дополнить отечественное законодательство положениями о том, что в любых ситуациях управленческие решения на основе анализа данных должен принимать исключительно человек;

- об использовании дипфейков и установлении ответственности за их неправомерное использование, включая гражданско-правовую, административную и уголовную ответственность. Помимо этого, требуется установление пределов использования дипфейков, как это было сделано в Китае и США. Например, 3 октября 2019 г. в штате Калифорния были опубликованы два законопроекта, подписанные губернатором. Assembly Bill № 730, запрещающий публикацию в средствах массовой информации, вводящих в заблуждение аудио- или видеоматериалов с кандидатом на государственную должность, с намерением нанести ущерб репутации кандидата или обманом заставить избирателя проголосовать за или против кандидата. Законопроект определял «существенно вводящие в заблуждение аудио- или визуальные носители» как изображения или аудио- или видеозапись внешнего вида, речи или поведения кандидата, которыми намеренно манипулировали таким образом, что изображение,

аудио- или видеозапись ложно показались бы разумному лицу подлинными и привели бы к тому, что у разумного человека возникло бы принципиально иное понимание или впечатление о выразительном содержании изображения или аудио- или видеозаписи, чем если бы он слышал или видел неизмененную, оригинальную версию изображения или аудио- или видеозаписи»¹⁹.

Второй законопроект, Assembly Bill № 602, запрещал публикацию контента порнографического характера, в том числе «измененных» видеороликов без согласия изображенного лица²⁰. Эти законопроекты критиковались как недостаточные для борьбы с дипфейками. Позднее в разных штатах было разработано еще несколько документов того же направления, готовятся и федеральные нормы (Postarnak, 2023:271).

В КНР, где интернет-пространство контролируется властями значительно плотнее, чем в западных странах, отказались от полного запрета дипфейков, однако любое использование дипфейков должно сопровождаться указанием на ложность предоставляемой информации. За несоблюдение этого требования предусмотрена уголовная ответственность²¹.

В Европейском Союзе идет процесс согласования норм о дипфейках, вырабатываемых национальными законодателями стран-членов ЕС. Важная роль в выработке таких норм принадлежит крупнейшим транснациональным социальным сетям, для которых ограничение дипфейков и вообще деятельности ИИ, – это неотложная практика (Postarnak, 2023:271–272). Их опыт обязательно должен быть учтен отечественными законодателями.

Мы считаем, что китайский вариант вполне рационален и применим в отечественных реалиях: в развлекательных и творческих целях дипфейки должны быть допустимы, однако видео обязательно должно быть промаркировано, чтобы зритель изначально понимал, что смотрит видео, сгенерированное нейросетью или с измененными элементами (вставками).

Предложенные нами нормы являются основными примерами возможных законодательных изменений, которые можно принять в рамках минимизации рисков использования искусственного интеллекта в сфере политики и управления. Однако они все же направлены на человека, т.е. подразумевают, что любой ИИ будет находиться под контролем конкретных людей. В то же время, в западной юридической науке обсуждается вопрос о придании искусственному интеллекту нового юридического статуса, вроде «цифрового юридического лица». В случае, если ИИ будет действовать в обществе полностью самостоятельно, соблюдение им правовых норм, скорее всего, не будет «человеческим». Эта проблема также требует изучения техническими специалистами, а также юристами, философами, психологами, и т.д. – проблема эта настолько важна для человечества будущего, что для ее изучения должны быть привлечены все необходимые ресурсы.

¹⁹ Assembly Bill № 730. Режим доступа: https://leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill_id=201920200AB730 (дата обращения: 01.08.2024).

²⁰ Assembly Bill № 602. Режим доступа: https://leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill_id=201920200AB602 (дата обращения: 01.08.2024).

²¹ Criminal Law of the People's Republic of China. Сайт Congressional-Executive Commission on China. Режим доступа: <https://www.cecc.gov/resources/legal-provisions/criminal-law-of-the-peoples-republic-of-china> (дата обращения: 04.11.2024).

Заключение

Технологии ИИ, несомненно, изменяют существующее законодательство, и, возможно, приведут к возникновению совершенно новых областей права. Эти технологии развиваются в последние годы «не по дням, а по часам», что требует от законодателей, политиков и юристов быстрой и адекватной реакции. Изменения, вносимые искусственным интеллектом во все стороны нашей жизни, настолько велики, что многие исследователи говорят о том, что происходит новый этап научно-технической революции. Столь сильные инструменты, применимые практически в любой деятельности, приносят не только новые возможности, но значительные риски для общества и всей планеты. Эти риски весьма разнообразны: военные, экологические, коренная перестройка рынка труда и структуры занятости населения, ошибки искусственного интеллекта, прямо противоправная деятельность людей с использованием новых технологий, законная деятельность, несущая неявную опасность для граждан и государства, и т.д.

Все эти новые вызовы обязательно должны быть отражены в действующем законодательстве, как на национальном, так и на международном уровнях. В данной статье мы предложили некоторые наиболее необходимые меры. Должны быть созданы новые структуры для контроля над новыми технологиями и соблюдением касающегося этой сферы законодательства.

Практика показывает, что в развитых странах уже идут все эти процессы, но идут по-разному: в зависимости от государственного устройства, политического режима, уровня развития применяемых технологий и т.д. Считаем важным изучать опыт всех участников этого процесса, включая и крупные корпорации, и применять тот опыт, который максимально подойдет под отечественные условия. В этой сфере вряд ли стоит устанавливать правила раз и навсегда. Вместо этого законодатели должны быть готовы быстро корректировать их с учетом эффективности принятых мер, и эволюции технологий. Считаем необходимым создание универсальных международных этических правил использования ИИ, в центре которых будут стоять интересы не только государств, но, в первую очередь людей, уже сегодня взаимодействующих с этими технологиями ежедневно.

References / Список литературы

- Basheva, O. A. (2020) Digital activism as a new method of civil mobilization. *Scientific result of Sociology and management*. 6 (1), 41–57. (in Russian).
Башиева О.А. Цифровой активизм как новый метод гражданской мобилизации // Научный результат. Социология и управление. 2020. Т. 6, № 1. С. 41–57.
- Bertovsky, L.V. (2021) High-tech law: Understanding, genesis and prospects. *RUDN Journal of Law*. 25 (4), 735–749. (in Russian).
Бертовский Л.В. Высокотехнологичное право: понятие, генезис и перспективы // Вестник Российского университета дружбы народов. Серия: Юридические науки. 2021. Т. 25, № 4. С. 735–749.
- Chernogor, N.N. (2022) Artificial intelligence and its role in the transformation of modern law. *Journal of Russian Law*. 26 (4), 5–15. (in Russian).
Черногор Н. Н. Искусственный интеллект и его роль в трансформации современного правового порядка // Журнал российского права. 2022. Т. 26, № 4. С. 5–15.
- Crawford, K. & Calo, R. (2016) There are gaps in artificial intelligence research. *Nature*. 538, 311–313.

- Gavrilova, Yu. A. (2021) The concept of integrating artificial intelligence into the legal system. *RUDN Journal of Law*. 25 (3), 673–692. <https://doi.org/10.22363/2313-2337-2021-25-3-673-692>
- Ilina, E.M. (2022) Politics and management in the Government of the Russian Federation: The political course of public administration. *Ars Administrandi*. 3, 403–421. (in Russian).
Ильина Е.М. Политика и управление в условиях цифровой трансформации: политологический ракурс искусственного интеллекта // *Ars Administrandi*. 2022. № 3. С. 403–421.
- Kazaryan, V.P. & Shutova, K.Yu. (2023) Ethical risks in the practice of artificial intelligence. *Ideas and ideals*. (2–2), 346–364. (in Russian).
Казарян В.П., Шутова К.Ю. Этические риски в практике искусственного интеллекта // *Идеи и идеалы*. 2023. № 2–2. С.346–364.
- Kiselev, A.S. (2021) On the need for legal regulation in the field of artificial intelligence: Deepfake as a threat to national security. *Bulletin of the Moscow State University. Series: Jurisprudence*. (3), 54–64. (in Russian).
Киселев А.С. О необходимости правового регулирования в сфере искусственного интеллекта: дипфейк как угроза национальной безопасности // *Вестник МГОУ. Серия: Юриспруденция*. 2021. № 3. С. 54–64.
- Khabrieva, T. Ya. (2023) Technological imperatives of the modern world and law. *Journal of Foreign Legislation and Comparative Jurisprudence*. 19 (1), 5–12. (in Russian).
Хабриева Т.Я. Технологические императивы современного мира и право // *Журнал зарубежного законодательства и сравнительного правоведения*. 2023. Т. 19. № 1. С. 5–12.
- Kozlov, S.E. (2020) Internet activism as a form of political participation in modern Russia. *State administration Electronic Bulletin*. (79), 154–169. <https://doi.org/10.24411/2070-1381-2020-10053> (in Russian).
Козлов С.Е. Интернет-активизм как форма политического участия в современной России // *Государственное управление. Электронный вестник*. 2020. № 79. С. 154–169. <https://doi.org/10.24411/2070-1381-2020-10053>
- Lemaikina, S.V. (2023) Problems of countering the use of deepfakes for prestigious purposes. *Lawyer-Jurist*. (2 (105)), 143–148. (in Russian).
Лемайкина С.В. Проблемы противодействия использования дипфейков в преступных целях // *ЮП*. 2023. № 2 (105). С. 143–148.
- Lovink, G. (2019) *Critical theory of the Internet*. Moscow, Ad Marginem Publ. (in Russian).
Ловинк Г. Критическая теория интернета. М. : Ad Marginem, 2019. 304 с.
- Novitsky, N.A. (2020) Issues of combining public administration in general, reducing investment risks in the development of digital systems with artificial intelligence. *Eurasian Union of Scientists*. (8–3 (77)), 48–54. (in Russian).
Новицкий Н.А. Вопросы совершенствования государственного управления в целях снижения инвестиционных рисков при развитии цифровых систем с искусственным интеллектом // *Евразийский Союз Ученых*. 2020. № 8–3 (77). С. 48–54.
- Postarnak, A.M. (2023) Analysis of foreign legal regulation of deepfakes: The problem of intellectual integrity protection. *Theory and practice of social development*. (6 (182)), 269–273. (in Russian).
Постарнак А.М. Анализ зарубежного правового регулирования дипфейков: проблема защиты интеллектуальной собственности // *Теория и практика общественного развития*. 2023. № 6 (182). С. 269–273.
- Sorokova, E.D. (2022) Ethics of artificial intelligence in Russia and the European Union: Common in risk-oriented approaches. *Discourse-Pi*. (3), 157–169. (in Russian).
Сорокова Е.Д. Этика искусственного интеллекта в России и Европейском союзе: общее в риск-ориентированных подходах // *Дискурс-Пи*. 2022. № 3. С. 157–169.
- Ren, Hu. (2022) Artificial intelligence and the legal system of the future: The right to access the electric energy of artificial intelligence. *Open Journal of Social Sciences*. 10 (12), 447–458.

- Volodenkov, S.V. & Fedorchenko, S.N. (2021) Digital infrastructures of civil and political activism: Current challenges, risks and limitations. *Monitoring*. 6, 97–118. (in Russian).
Володенков С.В., Федорченко С.Н. Цифровые инфраструктуры гражданско-политического активизма: актуальные вызовы, риски и ограничения // Мониторинг. 2021. № 6. С. 97–118.
- Volodenkov, S.V. (2022) Politics in digital format in the research of Russian and foreign scientists: Introducing the issue. *RUDN Journal of Political Science*. 24 (3), 339–350. <https://doi.org/10.22363/2313-1438-2022-24-3-339-350> (in Russian).
Володенков С.В. Политика в цифровом формате в исследованиях российских и зарубежных ученых: представляю номер // Вестник РУДН. Серия: Политология. 2022. № 3. С. 339–350. <https://doi.org/10.22363/2313-1438-2022-24-3-339-350>
- Yastrebov, O.A. (2018) Artificial intelligence in the legal space. *RUDN Journal of Law*. 22 (3), 315–328. (in Russian).
Ястребов О.А. Искусственный интеллект в правовом пространстве // Вестник Российской Федерации дружбы народов. Серия: Юридические науки. 2018. Т. 22. № 3. С. 315–328.
- Yuste, R., Goering, S. et al. (2017) These four priorities for neurotechnology and artificial intelligence. *Nature*. 551 (7679), 159–163.

Сведения об авторе:

Киселев Александр Сергеевич – кандидат юридических наук, ведущий научный сотрудник, Центр исследований и экспертиз; доцент кафедры международного и публичного права, юридический факультет, ФГБОУ ВО «Финансовый университет при Правительстве Российской Федерации»; заместитель главного редактора журнала «Банковское право» ИГ «Юрист»; 125468, Российская Федерация, г. Москва, Ленинградский проспект, д. 49.

ORCID: 0000-0002-4721, SPIN-код: 3570-8911

e-mail: alskiselev@fa.ru

About the author:

Alexander S. Kiselev – Candidate of Legal Sciences, Leading Researcher at the Center for Research and Expertise, Associate Professor of the Department of International and Public Law at the Faculty of Law of the Financial University under the Government of the Russian Federation, Deputy editor-in-chief of the Banking Law journal, Lawyer Publishing group; 49 Leningradsky Prospekt str., Moscow, 125468, Russian Federation

ORCID: 0000-0002-5044-4721, SPIN-code: 3570-8911

e-mail: alskiselev@fa.ru