

Ю. С. Харитонов

Московский государственный университет имени М.В. Ломоносова
г. Москва, Российская Федерация

Правовые подходы к регулированию искусственного интеллекта: уроки китайского права

Резюме. Принятие в ЕС Закона об искусственном интеллекте и включение в план работы законодательных органов Китая разработки соответствующего единого закона об ИИ показывает, что от этических кодексов и стратегий развития технологии мир переходит к вопросу всеохватного регулирования применения ИИ. В российской литературе на фоне большого внимания к актам ЕС недостаточное внимание уделено регулированию ИИ в Китае. Анализ показал, что уникальный опыт сегментарного регулирования в Китае и в определенной степени экспериментально-правового подхода в России позволяет действовать максимально аккуратно в создании системы регулятивной песочницы, формируя достаточно узко дифференцированное регулирование систем ИИ. Китайский подход к соблюдению социалистической морали и нравственности при применении ИИ в духе «развития науки и технологий во благо» и «подхода, ориентированного на человека» соответствует тенденциям российского права к защите традиционных духовно-нравственных ценностей. Китайская схема регулирования рисков классифицирует риски не горизонтально, а в зависимости от возникающих из одного и того же продукта, что больше подходит для таких технологий, как генеративный ИИ. Это может быть учтено при формировании российской позиции к содержанию правового регулирования применения ИИ.

Ключевые слова: искусственный интеллект; генеративный искусственный интеллект; цифровые технологии; китайское право; цифровое право; прозрачность искусственного интеллекта; объяснимость искусственного интеллекта; классификация систем искусственного интеллекта

Для цитирования: Харитонов Ю. С. Правовые подходы к регулированию искусственного интеллекта: уроки китайского права. *Lex russica*. 2025. Т. 78. № 5. С. 130–138. DOI: 10.17803/1729-5920.2025.222.5.130-138

Legal Approaches to Artificial Intelligence Regulation: Chinese Law Lessons

Yulia S. Kharitonova

Lomonosov Moscow State University
Moscow, Russian Federation

Abstract. The adoption of the Law on Artificial Intelligence in the EU and the inclusion in the work plan of the Chinese legislative bodies of the development of a corresponding unified law on AI shows that the world is moving from ethical codes and technology development strategies to the issue of comprehensive regulation of the use of AI. In the Russian literature, against the background of great attention to the EU acts, the regulation of AI in China receives insufficient attention. The analysis showed that the unique experience of segmental regulation in China and, to a certain extent, the experimental legal approach in Russia allows us to act as carefully as possible in creating a regulatory sandbox system, forming a rather narrowly differentiated regulation of AI systems. The Chinese approach to observing socialist morality and ethics in the application of AI in the spirit of «developing science and technology for the benefit» and «a human-centered approach» corresponds to the trends of Russian

law towards the protection of traditional spiritual and moral values. The Chinese risk management scheme classifies risks not horizontally, but depending on those arising from the same product, which is more suitable for technologies such as generative AI. This can be taken into account when forming the Russian position on the content of legal regulation of the use of AI.

Keywords: artificial intelligence; generative artificial intelligence; digital technologies; Chinese law; digital law; transparency of artificial intelligence; explainability of artificial intelligence; classification of artificial intelligence systems

Cite as: Kharitonova YuS. Legal Approaches to Artificial Intelligence Regulation: Chinese Law Lessons. *Lex Russica*. 2025;78(5):130-138. (In Russ.). DOI: 10.17803/1729-5920.2025.222.5.130-138

Введение

Развитие технологий искусственного интеллекта (ИИ) привело к возникновению значимых вызовов перед правовыми порядками всего мира. При этом одним из ключевых вопросов сегодня является обеспечение прозрачности и объяснимости ИИ на фоне недопущения нарушения фундаментальных прав граждан и обеспечения информационной безопасности в Сети. В настоящее время принципы разработки технологий ИИ рассматриваются на международном и национальном уровнях. Деятельность по надзору и управлению в области искусственного интеллекта становится всё более интенсивной во всем мире.

Ключевым актом в России является Национальная стратегия развития искусственного интеллекта на период до 2030 года¹, в которой выделены основные принципы развития и использования технологий ИИ. Выделенные в российском законодательстве принципы регулирования разработки технологий не получили глубокого толкования в правоприменительной практике или разъяснениях уполномоченных органов и в большой степени носят рекомендательный характер. При этом перечисленные в Стратегии принципы частично перекликаются с наборами аналогичных постулатов, отраженных в этических кодексах и рекомендациях в России и за рубежом.

Цели и последствия правового регулирования в Китае и ЕС

Одним из наиболее ярких примеров зарубежного опыта регулирования принципов применения технологии ИИ, безусловно, можно назвать европейский Закон об искусственном интеллекте² (AIA). Этот Закон призван ввести первое всеобъемлющее международное регулирование систем искусственного интеллекта, представляющее такие новшества, как риск-ориентированный подход к ИИ и акцент на защите фундаментальных прав человека, а также конституционных ценностей в сфере, где ИИ постоянно присутствует. Согласно ст. 3 AIA, системы ИИ представляют собой машинные системы, предназначенные для автономной работы, потенциально адаптирующиеся после развертывания и обрабатывающие входные данные для создания выходных результатов, таких как прогнозы, контент, рекомендации или решения, влияющие на физическую или виртуальную среду. Системы ИИ делятся на три категории:

1. АИС высокого риска: как правило, системы, используемые в критической инфраструктуре, здравоохранении и правоохранительных органах, которые должны соответствовать строгим требованиям безопасности, управления рисками и постоянного мониторинга в связи с их значительным влиянием на безопасность и основные права.

2. АИС с ограниченным риском: АИС, к которым применяются умеренные правила, требующие гарантий для предотвращения злоупотреблений и обеспечения прозрачности, что

¹ Указ Президента РФ от 10.10.2019 № 490 «О развитии искусственного интеллекта в Российской Федерации» (вместе с Национальной стратегией развития искусственного интеллекта на период до 2030 года) // СЗ РФ. 2019. № 41. Ст. 5700.

² COM/2021/206 Final. Proposal For A Regulation Of The European Parliament And Of The Council Laying Down Harmonised Rules On Artificial Intelligence (Artificial Intelligence Act) And Amending Certain Union Legislative Acts // URL: https://www.presidioeuropa.net/blog/wp-content/uploads/2024/02/Nuovo-Testo-TEN-T-2024-Annex-to-the-letter-of-9-February-2024_FINAL_EN.pdf.

позволяет устранить потенциальное воздействие без жестких мер, применяемых к системам с высоким риском.

3. АИС с минимальным риском: АИС, представляющие незначительные риски. Они в значительной степени освобождены от нормативного надзора, но всё же должны придерживаться основных принципов прозрачности, чтобы информировать пользователей о взаимодействии с ИИ.

Для систем всех видов предусмотрено обеспечение прозрачности и объяснимости работы ИИ. Такой подход получил закрепление после длительного периода консолидации мнений экспертов разных стран. Однако не является единственным путем регулирования в сфере ИИ.

Китайские исследователи подчеркивают, что применение ИИ не только повышает эффективность в отдельных областях человеческой деятельности, но и способствует развитию экономики и общества³. Некоторые исследователи утверждают, что влияние технологий на человеческое общество огромно и глубоко, учитывая идею, что «развитие технологий приводит к социальной эволюции»⁴. В то же время ученые исходят из осторожной позиции о том, что риски ИИ трудно предсказать на ранних стадиях технологического развития, хотя на современном этапе наша способность контролировать технологическое развитие относительно сильна. Но по мере развития технологии возможности контролировать технологии станут крайне ограниченными, потому что в это время «технология приобретет достаточную мощь и свой собственный путь развития»⁵.

Безопасность, неприкосновенность частной жизни и дискриминация составляют основу регулирования ИИ в Китае. При разработке и

использовании ИИ субъекты могут внедрять алгоритмическую предвзятость и предвзятость больших данных в ИИ в процессе построения задач, анализа данных и выбора критериев работы, что приводит к дискриминационным результатам. Этим задачам и посвящены конкретные нормы права. Принцип прозрачности ИИ рассматривается в китайском праве в непосредственной связи с объяснимостью и интерпретируемостью ИИ, которые называют связующим звеном между ИИ и правом. Поскольку «юридически оформить можно только то, что уже понятно», интерпретируемость ИИ стала необходимым условием для продвижения и применения ИИ, а также для решения проблемы его юридической ответственности⁶.

Опыт Китая в регулировании применения ИИ

Регулирование технологии искусственного интеллекта в Китае идет по собственному пути. Усилия законодателей и правоведов направлены на поиск способов минимизации рисков применения ИИ, а также поиска справедливого распределения ответственности.

В 2017 г. Государственный совет Китайской Народной Республики опубликовал План развития искусственного интеллекта следующего поколения⁷. План представляет собой акт национального стратегического уровня, в котором установлено, что ИИ является лидирующей технологией для будущего. Планируется, что к 2030 г. «теория, технология и применение ИИ в Китае в целом выйдут на мировой уровень, что сделает Китай крупным мировым центром инноваций в области ИИ и даст заметные результаты в развитии “умной” экономики, которая станет важной основой для превращения Китая

³ Яо З. В., Ли З. Л. Правовое регулирование рисков генеративного контента искусственного интеллекта // Журнал Сианьского университета Цзяотун (издание по общественным наукам). 2023. № 43 (05). С. 147.

⁴ Пу К., Сян В. Возможности и проблемы, возникающие при использовании ChatGPT в качестве генеративного ИИ, и стратегия реагирования // Журнал Чунцинского университета (издание по социальным наукам). 2023. № 29 (3). С. 102–114.

⁵ Шэнь Ф. Риски и управление генеративным искусственным интеллектом: решение «дилеммы Коллингриджа» // Журнал Чжэцзянского университета (издание по гуманитарным и социальным наукам). 2024. № 6. С. 73.

⁶ Лю Я. Исследование интерпретируемости искусственного интеллекта и юридической ответственности ИИ. 2022 // URL: <https://www.pkulaw.com/specialtopic/c85e5a72b1dd1e16fffd1d820bf85550bdfb.html?key word=%E4%BA%BA%E5%B7%A5%E6%99%BA%E8%83%BD%E7%9A%84%E5%8F%AF%E8%A7%A3%E9%87%8A%E6%80%A7%E5%8E%9F%E5%88%99&way=listView>.

⁷ Новое поколение искусственного интеллекта // URL: https://www.gov.cn/zhengce/content/2017-07/20/content_5211996.htm.

в ведущую инновационную и экономическую державу»⁸. В то же время изначально предполагалось, что к 2025 г. Китай только «устанавливает правовые нормы, этические стандарты и систему политики в области ИИ, формируя возможности оценки и контроля безопасности ИИ»⁹. П. В. Трошинский подчеркивает, что «проводимая китайскими властями политика в области законодательства характеризуется неспешной разработкой и последующим принятием необходимых для страны локальных актов правотворчества»¹⁰. Сейчас же в Китае всё чаще возникает дискуссия о необходимости обеспечения единства законодательства в данной сфере или даже о подготовке единого закона об ИИ, как в Евросоюзе. Закон об искусственном интеллекте был включен в план законодательной работы Госсовета КНР на 2023–2024 гг. Независимо от готовности государства принять единый акт, многие придерживаются мнения, что законодательство должно создать плюралистическую нормативную систему регулирования искусственного интеллекта, уточнить функции национальных и отраслевых стандартов, расширить пространство для групповых стандартов и сосредоточить внимание на роли этических норм¹¹.

Китайские юристы очень осторожны в разработке конечных правил и двигаются по этому направлению «малыми шагами». С 2022 г. Китай начал последовательно принимать нормативные документы, касающиеся алгоритмических рекомендаций, технологий глубокого синтеза и генеративных услуг ИИ. Регулирование ИИ строится исходя из областей его применения. Так, с 1 марта 2022 г. вступило в силу Положение об управлении алгоритмическими рекомендациями в информационных интернет-услугах¹² (далее — Положение об алгоритмах), согласно которому была введена норма о создании системы регистрации алгоритмов. Компании обязаны подавать заявки на регистрацию

алгоритмов через систему регистрации алгоритмов Управления киберпространства Китая (САС, 国家互联网信息办公室), раскрывая основные параметры работы технологии.

Положение об управлении глубоким синтезом в информационных интернет-сервисах¹³ от 25.11.2022 было принято для усиления всестороннего управления информационными услугами в Интернете и создало основу для регистрации систем ИИ. Процедуры и правила регистрации стандартизированы на основе практики. 10 июля 2023 г. Управление киберпространства Китая опубликовало Временные меры по управлению услугами генеративного искусственного интеллекта¹⁴, вступившие в силу 15 августа 2023 г. В данном акте определена концепция генеративного ИИ и обязательства для поставщиков соответствующих продуктов и услуг. Тем самым Китай внедрил механизм «двойной регистрации», состоящий из системы регистрации алгоритмов и системы регистрации ИИ (для больших языковых моделей). Регистрация больших языковых моделей требует наиболее эффективного взаимодействия регулирующих органов и поставщиков услуг для накопления опыта регулирования и разработки четких и конкретных правил, которые будут стимулировать предприятия выполнять свои обязательства по регистрации алгоритмов, особенно в части оценки их безопасности в области больших языковых моделей. Раскрытие алгоритмов при регистрации, по мнению контролирующих органов, обеспечит прозрачность применения ИИ в конкретных сферах.

Параллельно с национальными правилами и нормами вводились местные акты, такие как Положение о развитии индустрии искусственного интеллекта в специальной экономической зоне Шэньчжэнь (《深圳经济特区人工智能产业促进条例》) от 01.11.2022 или Положение о развитии индустрии искусственного интеллекта в Шанхае (《上海市促进人工智能产业发展条例》) от 10.10.2022.

⁸ URL: https://www.gov.cn/zhengce/content/2017-07/20/content_5211996.htm.

⁹ URL: https://www.gov.cn/zhengce/content/2017-07/20/content_5211996.htm.

¹⁰ Трошинский П. В. Развитие китайского законодательства в последние годы // Международное публичное и частное право. 2015. № 3. С. 36–40.

¹¹ Сун Х. Разработка нормативной структуры в законодательстве об искусственном интеллекте // Журнал Восточно-Китайского университета политических наук и права. 2024. № 5. С. 6.

¹² 互联网信息服务算法推荐管理规定 // URL: https://www.gov.cn/gongbao/content/2022/content_5682428.htm.

¹³ 互联网信息服务深度合成管理规定 // URL: https://www.gov.cn/zhengce/zhengceku/2022-12/12/content_5731431.htm.

¹⁴ 生成式人工智能服务管理暂行办法 // URL: https://www.gov.cn/zhengce/zhengceku/202307/content_6891752.htm.

То есть в решение вопросов применимости конкретных продуктов ИИ включалось и региональное нормотворчество.

В названных актах установлено, что развитие индустрии ИИ в городе следует принципам технологичности, ориентированности на приложения, ориентированности на людей, принципам безопасности и управляемости. Причем изначально к ИИ в китайском праве выстраивалось отношение не как к отдельному объекту, а как к части индустрии. В частности, в Шанхайском Положении определялось, что термин «индустрия искусственного интеллекта» относится к основным отраслям, таким как исследования, разработка и производство программных и аппаратных продуктов, системных приложений и интегрированных услуг, связанных с искусственным интеллектом, а также интеграция и применение искусственного интеллекта, исследования в сфере жизнеобеспечения людей, социального управления, экономического развития и других смежных областях (ст. 3).

1 марта 2024 г. Национальным комитетом Китая по стандартизации технологий кибербезопасности были официально опубликованы Основные требования безопасности для услуг генеративного искусственного интеллекта¹⁵. Данный акт детализирует требования к реализации соответствующих положений Временных мер по управлению услугами генеративного ИИ, таких как законность источников данных, безопасность контента и т.д., и предоставляет эффективные методы для поставщиков услуг генеративного ИИ для оценки безопасности. Это не только улучшает возможности предприятий в области безопасности услуг генеративного ИИ, но и предоставляет стандарты для регулирующих органов для оценки уровня безопасности конкретных услуг.

Регулирование ИИ на уровне подзаконных актов не должно противоречить законам КНР, которые составляют основу правового регулирования информационного и технологического развития страны, таким как: Закон КНР о кибербезопасности¹⁶, Закон КНР о безопасности данных¹⁷, Закон КНР о защите персональных данных¹⁸, Закон КНР о научно-техническом прогрессе¹⁹, а также Положение об оценке безопасности информационных интернет-сервисов с атрибутами общественного мнения или возможностями социальной мобилизации²⁰, Правила управления информационными сервисами публичных учетных записей интернет-пользователей²¹. Эти нормативные акты образуют систему управления алгоритмами в различных секторах, особенно в предоставлении генеративных услуг ИИ. В целом регулирование применения ИИ в китайском праве касается установления обязанностей нескольких типов: обязательства, связанные с механизмами надзора; обязательства, связанные с обучением алгоритмов; обязательства, связанные с управлением контентом; обязательства поставщиков услуг в отношении пользователей.

Основные особенности китайского подхода к регулированию ИИ

На фоне расширения сфер применения искусственного интеллекта в Китае активно проводится работа по созданию механизма правового обеспечения принципа прозрачности. Во многом обращение к принципу прозрачности связывают с его ключевым значением для решения проблемы «черного ящика» алгоритма²². Исследователи отмечают, что концепция прозрачности включает в себя две составляющие: формальную и предметную (по существу)²³.

¹⁵ 生成式人工智能服务安全基本要求 // URL: <https://www.tc260.org.cn/front/postDetail.html?id=20240301164054>.

¹⁶ 中华人民共和国科学技术进步法 // URL: https://www.gov.cn/xinwen/2021-12/25/content_5664471.htm.

¹⁷ 互联网信息服务算法推荐管理规定 // URL: https://www.cac.gov.cn/2022-01/04/c_1642894606364259.htm.

¹⁸ 中华人民共和国个人信息保护法 // URL: https://www.cac.gov.cn/2021-08/20/c_1631050028355286.htm?eqid=b7a9c7a1000acbc7000000026465ed77.

¹⁹ 中华人民共和国科学技术进步法 // URL: https://www.gov.cn/xinwen/2021-12/25/content_5664471.htm.

²⁰ 具有舆论属性或社会动员能力的互联网信息服务安全评估规定 // URL: https://www.cac.gov.cn/2018-11/15/c_1123716072.htm.

²¹ 互联网用户公众账号信息服务管理规定 // URL: https://www.gov.cn/xinwen/2021-01/24/content_5582219.htm.

²² Чжан Ю. О правовой реализации принципа прозрачности в сфере искусственного интеллекта // Юридический факультет Южно-Китайского нормального университета. Серия политико-правовых дискуссий. 2024. № 2. С. 124.

²³ Чжан Ю. Указ. соч.

Формальная прозрачность подразумевает раскрытие базовой информации об искусственном интеллекте, делающее развертывание и использование искусственного интеллекта несекретным. Предметная прозрачность тесно связана с интерпретируемостью искусственного интеллекта, она подчеркивает важность объяснения раскрываемой информации в понятной форме, что должно сделать соответствующую информацию действительно познаваемой. На первый взгляд, такой подход созвучен европейскому. Статья 4 Положения об алгоритмах предусматривает, что предоставление услуг алгоритмических рекомендаций должно следовать принципам справедливости и честности, открытости и прозрачности. Поставщики услуг должны сформулировать и обнародовать правила выработки рекомендаций (ст. 7), оптимизировать прозрачность и интерпретируемость правил поиска, сортировки, выбора, продвижения и показа контента (ст. 12), информировать пользователей о принципах, цели и намерениях применения рекомендательных сервисов, а также об основных механизмах работы алгоритмических рекомендательных услуг²⁴.

Специфика китайского подхода состоит в следующем. Прежде всего, в Китае многие исследования ИИ ведутся через призму доктрины верховенства права²⁵. В данном контексте логично общее указание на то, что при реализации принципа верховенства права следует тщательно учитывать принципы прозрачности и объяснимости: с учетом технических условий необходимо устанавливать различные обязательства по раскрытию и разъяснению содержания информации в соответствии с возможностями регулирующих органов и ответственности.

Второй особенностью китайского подхода можно назвать стремление к классификации систем ИИ. В литературе отмечается, что в сфере цифрового права в Китае произошла смена парадигмы: от первоначального индивидуального предотвращения рисков юристы перешли к систематическому управлению рисками, основанному на иерархическом и классифи-

цированном управлении, управлении полным жизненным циклом ИИ-систем²⁶. Такой взгляд может быть сопоставлен с европейским. Как и в ЕС, классификация алгоритмов и управление безопасностью рассматриваются как основной инструмент управления безопасностью алгоритмов в будущем. Но и здесь мы обнаруживаем отличительные черты китайского подхода.

В Положении об алгоритмах установлены общие критерии классификации поставщиков услуг рекомендательных алгоритмов исходя из атрибутов общественного мнения или способности к социальной мобилизации сервиса, категории контента, количества пользователей, степени важности данных, обрабатываемых технологией рекомендаций алгоритмов, степени вмешательства в поведение пользователя и т.д. Положение об алгоритмах на уровне надзора за платформами предусматривает иерархическую и классифицированную систему управления, а для платформ общего назначения — обязанность вести веб-журналы и сотрудничать с государством. Для платформ, являющихся лидерами общественного мнения или обладающих возможностями социальной мобилизации, вводится система регистрации и требуется проактивная оценка безопасности. Те, кто предоставляет новостные и информационные услуги в Интернете, должны получить лицензию на предоставление таких услуг в соответствии с законодательством. С точки зрения применения Положения об алгоритмах эталонный стандарт классификации алгоритмов зависит и от степени «чувствительности» данных, но в настоящее время меняется на показатель степени «важности» данных, что в наибольшей степени соответствует нормам Закона о безопасности данных.

Обязанность поставщиков услуг принимать превентивные меры должна быть более строгой в случаях незаконного и нежелательного контента, связанного с национальной и общественной безопасностью, характеризующегося высоким риском, и относительно менее строгой в случаях незаконного и нежелательного кон-

²⁴ Харитонов Ю. С., Тяньфан Я. Рекомендательные системы цифровых платформ в Китае: правовые подходы и практика обеспечения прозрачности алгоритмов // Закон. 2022. № 9. С. 40–49.

²⁵ Чжан Ю. Указ. соч. ; Чжан С. Продвижение верховенства закона с главным направлением продвижения нового качества продуктивности цифрового интеллекта // Цифровое правовое государство. 2024. № 4. С. 1.

²⁶ Чжао Ц., Чжоу Ж. О смене парадигмы цифровых юридических исследований: системное управление рисками // Академический журнал Цюши. 2024. № 4. С. 136.

тента, связанного только с нарушением частных прав, характеризующегося низким риском²⁷.

Таким образом, для разных уровней риска ИИ в Китае принимаются стандарты реализации, основанные на классификации и сценариях, что способствует реализации принципа прозрачности в двух аспектах: защита пользователей и права общественности на информацию и обеспечение национальных регулирующих полномочий. Однако, в отличие от европейского подхода, в Китае придерживаются позиции, что с учетом технических условий необходимо дифференцировать требования к содержанию информации, подлежащей раскрытию и разъяснению, в соответствии с различными возможностями регулирующих органов и общественности, а также принять градуированные и классифицированные подсценарии внедрения стандартов для различных уровней риска различных ИИ.

Третьей значимой особенностью китайского подхода можно назвать приверженность требований к применению ИИ в рамках социалистической морали и нравственности. Общеизвестно, что дискриминация не должна создавать алгоритмическую власть и не должна нарушать права личности и права на личную информацию других лиц. В то же время особенностью китайского подхода является то, что, например, Временные меры по управлению услугами генеративного ИИ требуют, чтобы при предоставлении и использовании услуг генеративного ИИ соблюдались основные социалистические ценности, принимались эффективные меры по предотвращению дискриминации по признаку этнической принадлежности, убеждений, страны происхождения и т.д.; эти услуги не должны приводить к формированию алгоритмической власти или нарушать права личности, права и интересы других лиц в отношении личной информации, должны повышать прозрачность услуг генеративного искусственного интеллекта, а также точность и надежность генерируемого контента.

Положение об алгоритмах устанавливает стандарты информационных услуг в ст. 6–10, которые в основном посвящены контролю за незаконной и нежелательной информацией; так, в ст. 8 четко указано, что «никакая алгорит-

мическая модель не должна быть создана для того, чтобы побуждать пользователей к зависимости или чрезмерному потреблению информации в нарушение законов и правил или против этики и морали». Положение об алгоритмах также запрещает использование алгоритмов для манипулирования информацией и влияния на общественное мнение в Интернете. Положение об алгоритмах, таким образом, требует соблюдения не только писаных законов и правил, но и общественной морали и этики, деловой и профессиональной этики, принципов справедливости и честности, открытости и прозрачности, научной обоснованности, а также честности и благонадежности.

Нарушение основных социалистических ценностей и дискриминационный контент относятся к категории высокого риска и требуют особого внимания и управления. Развитие ИИ в Китае можно обеспечить путем создания и совершенствования системы тестирования, системы надзора и системы юридической ответственности²⁸. Определения рисков безопасности помогают поставщикам услуг лучше выявлять и избегать потенциальных административных нарушений.

Вывод

Обобщение опыта правового регулирования применения ИИ в Китае в сравнении с подходами Европейского Союза позволяет выявить некоторые идеи, которые могут быть восприняты в российской законотворческой практике. Прежде всего, и китайский, и европейский опыт показал, что единое законодательство об ИИ тяготеет к концепции всеохватности на фоне осмотристельности и градации классификации рисков. В то же время уникальный опыт сегментарного регулирования в Китае и в определенной степени экспериментально-правового подхода в России позволяет действовать максимально аккуратно при формировании системы регулятивной пещочницы, создавая достаточно узко дифференцированное регулирование систем ИИ. Во всех случаях вводится система контроля и надзора за соблюдением требований закона в отношении применения ИИ. Однако компетенция и органы

²⁷ Яо З. В., Ли З. Л. Правовое регулирование рисков генеративного контента искусственного интеллекта // Журнал Сианьского университета Цзяотун (издание по общественным наукам). 2023. № 43 (05). С. 147.

²⁸ Ян Ц. Развитие заслуживающего доверия искусственного интеллекта и построение правовой системы // Восточное правоведение. 2024. № 4. С. 95.

либо комиссии, уполномоченные осуществлять такой контроль, действуют на основе разных принципов. В то же время более жесткий, чем в ЕС, китайский подход к соблюдению социалистической морали и нравственности при применении ИИ в духе «развития науки и технологий во благо» («科技向善»), соответствует тенденциям российского права к защите традиционных духовно-нравственных ценностей, культуры и исторической памяти.

Некоторые ученые отмечают, подход ЕС к ограничениям применения высокорисковых систем ИИ может привести к чрезмерному регули-

рованию и негативно повлиять на развитие отрасли, например генеративного ИИ или работы с биометрическими данными. Китайская схема регулирования рисков, по сути, является переходом от классификации рисков для различных продуктов к классификации различных рисков, возникающих из одного и того же продукта²⁹, что больше подходит для таких технологий, как генеративный ИИ. Таким образом, при формировании российского подхода к правовому регулированию применения ИИ следует учесть позиции разных стран для минимизации рисков и выработки адекватных правовых решений.

СПИСОК ЛИТЕРАТУРЫ

Лю Я. Исследование интерпретируемости искусственного интеллекта и юридической ответственности ИИ. 2022. [刘艳红.人工智能的可解释性与AI的法律责任问题研究 2022].

Пу К., Сян В. Возможности и проблемы, возникающие при использовании ChatGPT в качестве генеративного ИИ, и стратегия реагирования // Журнал Чунцинского университета (издание по социальным наукам). 2023. № 29 (3). С. 102–114 [蒲清平, 向往.生成式人工智能 — ChatGPT的变革影响、风险挑战及应对策略[J].重庆大学学报(社会科学版), 2023,29(03):103].

Сун Х. Разработка нормативной структуры в законодательстве об искусственном интеллекте // Журнал Восточно-Китайского университета политических наук и права. 2024. № 5 [宋华琳.人工智能立法中的规制结构设计 //《华东政法大学学报》2024. № 5. С. 6].

Трощинский П. В. Развитие китайского законодательства в последние годы // Международное публичное и частное право. 2015. № 3. С. 36–40.

Харитонов Ю. С., Тяньфан Я. Рекомендательные системы цифровых платформ в Китае: правовые подходы и практика обеспечения прозрачности алгоритмов // Закон. 2022. № 9. С. 40–49.

Чжан С. Продвижение верховенства закона с главным направлением продвижения нового качества продуктивности цифрового интеллекта // Цифровое правовое государство. 2024. № 4 [张效羽.以促进数智化新质生产力为主线推进法治建设.《数字法治》双月刊2024年第4期第1].

Чжан Ю. О правовой реализации принципа прозрачности в сфере искусственного интеллекта // Юридический факультет Южно-Китайского нормального университета. Серия политико-правовых дискуссий. 2024. № 2 [张永忠.论人工智能透明度原则的法治化实现 // 华南师范大学法学院. 2024.2.124.].

Чжао Ц., Чжоу Ж. О смене парадигмы цифровых юридических исследований: системное управление рисками // Академический журнал Цюши. 2024. № 4 [赵精武; 周瑞珏.论数字法学研究范式的转向:风险体系化治理 //《求是学刊》].

Шэнь Ф. Риски и управление генеративным искусственным интеллектом: решение дилеммы Коллингриджа // Журнал Чжэцзянского университета (издание по гуманитарным и социальным наукам). 2024. № 6. [沈芳君.生成式人工智能的风险与治理 — 兼论如何打破《科林格里奇困境》《浙江大学学报(人文社会科学版)》2024,6:73].

Яо З. В., Ли З. Л. Правовое регулирование рисков генеративного контента искусственного интеллекта // Журнал Сианьского университета Цзяотун (издание по общественным наукам). 2023. № 43 (05) [姚志伟,李卓霖.生成式人工智能内容风险的法律规制[J].西安交通大学学报(社会科学版), 2023,43(05):147].

Ян Ц. Развитие заслуживающего доверия искусственного интеллекта и построение правовой системы // Восточное правоведение. 2024. № 4 [杨建.可信人工智能发展与法律制度的构建 //《东方法学》. 2024. № 4. С. 95].

Arrieta A. B. et al. Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI // Information fusion. 2020. T. 58. P. 82–115.

²⁹ Arrieta A. B. et al. Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI // Information fusion. 2020. T. 58. C. 82–115.

REFERENCES

- Arrieta AB et al. Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information fusion*. 2020;58:82-115.
- Kharitonova YuS, Tianfang Ya. Recommender systems of digital platforms in China: Legal approaches and practices for ensuring transparency of algorithms. *Zakon*. 2022;9:40-49. (In Russ.).
- Lu Ya. A study of the interpretability of artificial intelligence and the legal responsibility of AI [刘艳红.人工智能的可解释性与AI的法律责任问题研究 2022]. 2022. (In Chinese).
- Pu K, Xiang V. The opportunities and challenges that arise when using ChatGPT as a generative AI, and a response strategy. *Journal of Chongqing University (Social Sciences Publication)* [蒲清平, 向往. 生成式人工智能 — ChatGPT的变革影响、风险挑战及应对策略[J]. 重庆大学学报(社会科学版), 2023, 29(03):103]. 2023;29(3):102-114. (In Chinese).
- Shen F. Risks and management of generative artificial intelligence: Solving the Collingridge dilemma. *Journal of Zhejiang University (publication on Humanities and social Sciences)* [沈芳君. 生成式人工智能的风险与治理 — 兼论如何打破«科林格里奇困境»《浙江大学学报(人文社会科学版)》]. 2024;6:73. (In Chinese).
- Sun H. Development of a regulatory structure in the legislation on artificial intelligence. *Journal of the East China University of Political Sciences and Law* [宋华琳. 人工智能立法中的规制结构设计 //《华东政法大学学报》]. 2024;5:6. (In Chinese).
- Troshchinsky PV. The development of the Chinese legislation in recent years. *Mezhdunarodnoe publichnoe i chastnoe pravo*. 2015;3:36-40. (In Russ.).
- Yang Ts. The development of trustworthy artificial intelligence and the construction of a legal system. *Eastern Jurisprudence*. [杨建. 可信人工智能发展与法律制度的构建 //《东方法学》]. 2024;4:95. (In Chinese).
- Yao ZV, Lee ZL. Legal regulation of the risks of generative artificial intelligence content. *Journal of Xi'an Jiaotong University (Social Sciences publication)*. 2023;43(05):147. [姚志伟, 李卓霖. 生成式人工智能内容风险的法律规制[J]. 西安交通大学学报(社会科学版), 2023, 43(05):147] (In Chinese).
- Zhang S. Promoting the rule of law with the focus on promotion of a new quality of productivity of digital intelligence. *Digital Rule of law* [张效羽. 以促进数智化新质生产力为主线推进法治建设.《数字法治》双月刊2024年第4期第1]. 2024;4. (In Chinese).
- Zhang Yu. On the legal implementation of the principle of transparency in the field of artificial intelligence. Faculty of Law of South China Normal University. *A series of political and legal discussions*. [张永忠. 论人工智能透明度原则的法治化实现 // 华南师范大学法学院. 2024.2.124.]. 2024;2. (In Chinese).
- Zhao C, Zhou J. About the paradigm shift in digital legal research: Systemic risk management. *Academic Journal of Qiushi* [赵精武; 周瑞珏. 论数字法学研究范式的转向: 风险体系化治理 //《求是学刊》]. 2024;4. (In Chinese).

ИНФОРМАЦИЯ ОБ АВТОРЕ

Харитоновна Юлия Сергеевна, доктор юридических наук, профессор, руководитель НОЦ «Центр правовых исследований искусственного интеллекта и цифровой экономики», профессор кафедры предпринимательского права Московского государственного университета имени М.В.Ломоносова д. 1, Ленинские горы, г. Москва 119991, Российская Федерация
sovet2009@rambler.ru

INFORMATION ABOUT THE AUTHOR

Yulia S. Kharitonova, Dr. Sci. (Law), Professor, Head of the Research Center for Legal Research of Artificial Intelligence and Digital Economy, Professor, Department of Business Law, Lomonosov Moscow State University, Moscow, Russian Federation
sovet2009@rambler.ru

Материал поступил в редакцию 3 октября 2024 г.
Статья получена после рецензирования 31 октября 2024 г.
Принята к печати 15 апреля 2025 г.

Received 03.10.2024.
Revised 31.10.2024.
Accepted 15.04.2025.