

Научная статья

УДК: 340.13

DOI:10.17323/2072-8166.2023.3.245.267

Нормативно-правовое регулирование генеративного искусственного интеллекта в Великобритании, США, Европейском союзе и Китае



Ли Яо

Институт сравнительного правоведения при Китайском политико-правовом университете, Китайская Народная Республика 100088, Пекин, шоссе Ситу-чэн, 25, yaozaihenmang@mail.ru, <https://orcid.org/0000-0002-9933-7513>



Аннотация

30 ноября 2022 года компания OpenAI запустила разговорный искусственный интеллект ChatGPT. С обновлением последней версии ChatGPT продемонстрировал впечатляющую способность понимать естественный язык, что сделало его привлекательным инструментом для компаний и частных лиц, желающих обеспечить обслуживание и поддержку клиентов. В качестве обучающих данных для GPT-3 используются текстовые данные, в основном, из общедоступной информации в Интернете. GPT-4, в свою очередь, помимо текстовых данных, использует для обучения большое количество изображений и таким образом может обрабатывать как текстовые, так и графические данные. Появление генеративного искусственного интеллекта (ИИ) сильно повлияло на жизнь человека, однако нужно отметить, что интеллектуальные технологии — это обоюдоострый меч: быстрое развитие технологии генеративного ИИ, с одной стороны, может повысить эффективность и производительность, снизить затраты и открыть новые возможности для роста экономики. С другой, использование сервисов генеративного ИИ для создания синтетического контента в виде текста, аудио, видео и изображений создает возможные риски. На сегодняшний день разные регионы по всему миру находятся на разных стадиях разработки нормативных актов, касающихся генеративного ИИ. С использованием сравнительно-правового метода и метода системного анализа в настоящей статье подробно анализируются основные модели правового регулирования генеративного ИИ в современном мире на примере Великобритании, США, Евросоюза и Китая, отмечается разница в подходах при разработке и принятии соответствующих нормативных правовых актов в области регулирования

сервисов генеративного ИИ. Раскрываются позиции китайского правительства в отношении «развития и безопасности», и также «инновации и управления» на настоящий момент. Описываются основные тенденции совершенствования регулирования сервисов генеративного ИИ со стороны законодательства Китая. Сделан вывод о необходимости сбалансировать «верховенство закона» и «инновацию», и способствовать здоровому развитию генеративного ИИ.



Ключевые слова

генеративный искусственный интеллект; Chat Generative Pre-Trained Transformer; риски; правовое регулирование; инновации; зарубежный опыт.

Для цитирования: Ли Яо. Особенности нормативно-правового регулирования генеративного искусственного интеллекта в Великобритании, США, Евросоюзе и Китае // Право. Журнал Высшей школы экономики. 2023. Том 16. № 3. С. 245–267. DOI:10.17323/2072-8166.2023.3.245.267.

Research article

Specifics of Regulatory and Legal Regulation of Generative Artificial Intelligence in the UK, USA, EU and China



Li Yao

Institute of Comparative Law, China University of Political Science and Law, 25, Xitucheng Road, Haidian District, Beijing, China, yaozaihenmang@mail.ru, <https://orcid.org/0000-0002-9933-7513>



Abstract

On November 30, 2022, OpenAI launched ChatGPT conversational artificial intelligence; with the latest version update, ChatGPT has demonstrated an impressive ability to understand natural language, making it an attractive tool for companies and individuals looking to provide customer service and support. GPT-3 uses textual data, mostly from publicly available information on the Internet, as training data. GPT-4, on the other hand, uses a large number of images in addition to textual data for training, and thus can process both textual and graphical data. The emergence of generative AI has greatly impacted human life, but it can be said that intelligent technology is a double-edged sword: rapid development of generative artificial intelligence (AI) technologies, on the one hand, it can improve efficiency and productivity, reduce costs and open new opportunities for economic growth, but on the other hand, the use of generative AI services to create synthetic content in the form of text, audio, video and images poses possible risks. To date, different regions around the world are at different stages of development of normative acts concerning generative AI. Using the comparative legal method and the method of system analysis, this article analyzes in detail the main models of legal regulation of generative AI in the modern world on the example of the UK, the USA, the European Union and China, noting the different approach in the development and adoption of relevant normative legal

acts in the field of regulation of generative AI services. It especially reveals the Chinese government's position on "development and security" and "innovation and governance" at present. The main trends of improving the regulation of generative AI services by China are proposed, and it is concluded that it is necessary to balance "rule of law" and "innovation" and promote the healthy development of generative AI.



Keywords

Generative AI; Chat Generative Pre-Trained Transformer, risks; regulation; innovations; foreign experience.

For citation: Li Yao (2023) Specifics of Regulatory and Legal Regulation of Generative Artificial Intelligence in the UK, USA, EU and China. *Law. Journal of the Higher School of Economics*, vol. 16, no. 3, pp. 245–267 (in Russ.). DOI: 10.17323/2072-8166.2023.3.245.267.

Введение

Примерно в 2010 году появилась технология «глубокого обучения» (Deep Learning), которая позволила значительно улучшить способность компьютеров распознавать изображения, обрабатывать звук и т.д. 30 ноября 2022 года компания OpenAI запустила разговорный искусственный интеллект ChatGPT; с обновлением последней версии ChatGPT продемонстрировал впечатляющую способность понимать естественный язык, что сделало его привлекательным инструментом для компаний и частных лиц, желающих усовершенствовать обслуживание клиентов и тем самым упрочить их поддержку. По словам китайского специалиста, «ChatGPT — это важный прорыв в области искусственного интеллекта после поражения чемпиона Европы по игре в го от AlphaGo, ознаменовавший становление парадигмы интеллектуальных вычислений с большими моделями»¹. Генеративный ИИ определяют как технологию, которая использует модели глубокого обучения для создания генеративных информационных материалов (текст, изображения, видео и пр.) в ответ на запрос человека. Появление генеративного ИИ сильно повлияло на жизнь человека, но, можно сказать, что интеллектуальные технологии — это обоюдоострый меч, который сталкивается с целым рядом проблем, помогая при этом модернизировать государственное управление [Цюэ Т., Лу Ц., 2021: 28]. В ответ на ряд рисков и вызовов многие страны выдвинули относительно систематизированные программы управления и контроля за ИИ. Учитывая, что в настоящее время генеративный ИИ все еще находится на предварительной стадии развития, объяснение

¹ Available at: URL: <http://www.zidonghua.com.cn/news/tech/47847.html> (дата обращения: 06.07.2023)

режима его работы, опыты продвижения инноваций, прогнозирование связанных с ним рисков и внимание к законодательной динамике развитых стран в отрасли генеративного искусственного интеллекта могут быть полезны в будущем совершенствовании законодательства и соответствующих правовых норм.

1. Режим работы генеративного ИИ и связанные с ним риски

В качестве обучающих данных для GPT-3 используются текстовые данные, почерпнутые в основном из общедоступной информации в Интернете. GPT-4, в свою очередь, помимо текстовых данных использует для обучения большое количество изображений, и, таким образом, может обрабатывать как текстовые, так и графические данные. Появление термина «генеративный ИИ» отражает важный шаг в процессе создания интеллекта. С технической точки зрения такая крупномасштабная модель фактически выполняет задачу «предсказания следующего слова»: модель генерирует первое слово на основе «слова-подсказки», затем вводит в модель первое слово для генерации второго слова, затем вводит в модель первые два слова для генерации третьего слова и т.д., пока не будет произведен весь ответ; такой процесс называется «авторегрессия» (Autoregression)². В деталях это означает, что при ответе на вопрос пользователя система не извлекает и не обращается к существующим данным из базы данных или из Интернета (традиционная поисковая система), а предсказывает ответ, основываясь, в основном, на вероятности взаимосвязей между лингвистическими символами.

Таким образом, когда модель данных велика, она обладает возможностями, которых нет у обычных небольших моделей, например, рассуждениями на основе здравого смысла и написанием компьютерных программ. По этой логике если сделать модель больше, то она станет более практичным, гибким и универсальным инструментом, что полностью основано на эмерджентных способностях (Emergent Abilities) генеративного ИИ. Такая способность означает, что крупномасштабная модель может обладать неограниченным потенциалом [Wei J., Tay Y. et al., 2022:1-2]. GPT-2 имела 1.5 млрд. параметров, GPT-3 обладает 175 млрд. Хотя OpenAI решили пока не раскрывать количество параметров GPT-4 (помня о здравом смысле и учитывая тенденции развития), можно допустить, что GPT-4 обладает десятками триллионов парамет-

² Available at: URL: <https://mark-riedl.medium.com/a-very-gentle-introduction-to-large-language-models-without-the-hype-5f67941fa59e> (дата обращения: 19.07.2023)

ров³. Модель с большим количеством параметров сможет написать гораздо больше символов текста. В настоящее время в ChatGPT открыты два плагина: Browsing и Code Interpreter, которые позволяют составлять графики, карты, визуализировать данные, анализировать музыкальные плейлисты, создавать интерактивные HTML-файлы, отбирать наборы данных и извлекать цветовые палитры из изображений; также они легко взаимодействуют с более чем 5 000 приложений и могут обеспечить удобные решения в бронировании авиабилетов, гостиниц, в доставке товаров, в онлайн-шопинге, в юридической помощи и т.д.⁴

Хотя столь мощная эмерджентная способность открывает перед человечеством возможности развития, она несет в себе и ряд рисков. Израильский историк Юваль Ной Харари утверждает, что, манипулируя языком и генерируя его, генеративный ИИ вторгся в операционную систему человеческой цивилизации⁵.

Имеются примеры некоторых рисков: риски ложной информации: с одной стороны, сервис генеративного ИИ опирается на большой объем данных и информации для глубокого обучения, и эти данные в основном поступают из публичных текстовых данных в Интернете. Если сервис генеративного ИИ использует ложную информацию в качестве объекта обработки и анализа, результирующий текст также может содержать ложную информацию. Сервис генеративного ИИ обладает сильным автономным качеством, и даже если обрабатываемые данные получены из достоверных и точных источников, нельзя исключать генерации ложной информации в результате алгоритмической интеграции данных на основе сервиса генеративного ИИ. Еще риски заключаются в том, что в базовой структуре данных генеративного ИИ ныне преобладают англоязычные данные, и выходной контент неизбежно имеет иное понятие об истории и культуре неанглоязычных стран [Юань Ц., 2023: 23], таких как Китай и Россия, и порой генерирует неправдивую информацию. В этом случае население может подсознательно изменить свое долгосрочное понимание традиционной культуры и национальных особенностей после длительного получения ложной информации, содержащей ценностные предубеждения. Если генеративный ИИ будет использоваться в когнитивной войне на национальном уровне, то это создаст угрозу подрыва национального суверенитета [Чэнь Б., 2023: 14].

³ Available at: URL: <https://neuro-texter.ru/blog/vishel-gpt-4-razbor-vmeste-s-ekspertom?ysclid=lkph205t8306165304> (дата обращения: 18.07.2023)

⁴ Available at: URL: <https://openai.com/blog/chatgpt-plugins> (дата обращения: 17.07.2023)

⁵ Available at: URL: <https://www.economist.com/by-invitation/2023/04/28/yuval-noah-harari-argues-that-ai-has-hacked-the-operating-system-of-human-civilisation> (дата обращения: 17.07.2023)

Генеративный ИИ может быть связан с угрозами безопасности данных: сервис генеративного ИИ формирует модели, связываясь с большим массивом информации, включая релевантные данные, введенные пользователями. Как отмечают китайские ученые, «как только пользователь вводит соответствующие данные с помощью сервиса генеративного ИИ, они становятся частью процесса обучения машинного интеллекта, создавая риск безопасности для личной информации пользователя, конфиденциальной информации, коммерческих тайн и других секретных данных» [Цао Ш., Цао Р., 2023: 5]. 20.03.2023 уязвимость базы с открытым исходным кодом ChatGPT позволила некоторым пользователям видеть разговоры других пользователей, их имена, электронные адреса и даже платежную информацию⁶.

Кроме того, применение технологии генеративного ИИ может создать потенциальный риск для национальной безопасности данных, поскольку технология алгоритмов и заранее подготовленные наборы данных сервиса генеративного ИИ поставляются из-за рубежа. Ученые отмечают, что «ChatGPT также может быть объединен с другими данными для добычи и анализа с целью получения разведывательной информации, связанной с национальной безопасностью, общественной безопасностью и законными интересами частных лиц и организаций»⁷. Большая проблема на данный момент в том, что даже с учетом этих опасностей цифровые платформы избегают ответственности, например, OpenAI пытается избежать ответственности за нарушение данных, налагая на пользователей чрезмерные обязанности. В Условиях использования установлено, что «пользователь несет полную ответственность за контент»⁸. OpenAI не гарантирует надежности, пригодности, качества, соблюдения прав, точности выходных данных. Он даже «не несет ответственности за косвенные, случайные, специальные, последовательные или образцовые убытки, включая убытки от потери прибыли, деловой репутации, использования или потери данных, или другие потери, даже если известно о возможности таких убытков»⁹.

Многие страны осознают риски, которые несет человечеству генеративный ИИ, и делают все возможное, чтобы принять меры к их устране-

⁶ Available at: URL: <https://openai.com/blog/march-20-chatgpt-outage> (дата обращения: 28.06.2023)

⁷ Available at: URL: https://www.cssn.cn/fx/szfx/202304/t20230423_5624115.shtml (дата обращения: 26.06.2023)

⁸ OpenAI, Terms of use, 3.(a)(b)(c). Available at: URL: <https://openai.com/terms/> (дата обращения: 02.07.2023)

⁹ OpenAI, Terms of use, 7.(b)(c). Available at: URL: <https://openai.com/terms/> (дата обращения: 25.06.2023)

нию. 22.03.2023 глава Tesla и SpaceX Илон Маск, один из основателей Apple Стив Возняк и больше тысячи экспертов опубликовали открытое письмо. Авторы письма призвали разработчиков искусственного интеллекта немедленно остановить минимум на полгода обучение более мощных систем, чем GPT-4 от компании Open AI пока не появятся общие правила безопасности. По их мнению, если обучение не удастся остановить в ближайшее время, власти «должны вмешаться и ввести мораторий»¹⁰.

Американский Center for AI and Digital Policy (CAIDP) подал официальную жалобу на OpenAI в Федеральную торговую комиссию. Центр CAIDP призвал власти провести тщательную проверку компании на предмет «недобросовестной деловой практики»¹¹. 31.03.2023 власти Италии первыми в мире запретили доступ к чат-боту ChatGPT. Заявление об этом размещено на сайте итальянского Национального управления защиты персональных данных. В компании рассказали, что были вынуждены временно отключить ChatGPT из-за «ошибки в работе библиотеки с открытым исходным кодом», которая позволила некоторым пользователям видеть сообщения других лиц, общавшихся с чат-ботом¹². Но отключение временное, а более важная задача на данный момент заключается в разработке закона, чтобы избежать всех рисков. Рассмотрим некоторые примеры из мирового законодательного опыта.

2. Зарубежный опыт правового регулирования генеративного ИИ

На сегодняшний день разные регионы мира находятся на различных стадиях разработки нормативных актов, касающихся генеративного ИИ.

2.1. Великобритания и Соединенные Штаты Америки

Эти две страны в основном выступают за «мягкое регулирование», основной целью которой является стимулирование инноваций. 29.03.2023 правительство Великобритании опубликовало программный документ «Проинновационный подход к регулированию искусственного интеллекта» (A pro-innovation approach to AI regulation)¹³. В целом это прави-

¹⁰ Available at: URL: <https://futureoflife.org/open-letter/pause-giant-ai-experiments/> (дата обращения: 21.07.2023)

¹¹ Available at: URL: <https://incrypted.com/zhaloby-na-openai/> (дата обращения: 21.07.2023)

¹² Available at: URL: <https://www.politico.eu/article/chatgpt-italy-lift-ban-garanteprivacy-gdpr-openai/> (дата обращения: 21.07.2023)

¹³ A pro-innovation approach to AI regulation. Available at: URL: <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach/white-paper> (дата обращения: 21.07.2023)

тельство считает, что ввиду темпов развития технологий ИИ необходимо гибкое правовое регулирование. Введение жестких и обременительных законодательных требований для новых компаний может затормозить инновации в области ИИ и ограничить возможность быстрого реагирования на будущие технологические прорывы. Программные документы не следует применять огульно, необходим гибкий подход к реализации правовых принципов, опирающийся на сценарии действий отраслевых регулирующих органов.

США, исходя из необходимости собственного международного лидерства в области ИИ, также придерживаются относительно «мягкого регулирования». 11.02.2019 президент Дональд Трамп подписал приказ «Поддержание американского лидерства в области искусственного интеллекта»¹⁴. В приказе сформулированы общие направления управления ИИ, ориентированные на укрепление глобального лидерства. 04.02.2022 Управление научно-технической политики Белого дома опубликовало проект «Билля о правах ИИ» (Blueprint for an AI Bill of Rights)¹⁵. Документ включает пять принципов: создание безопасных и эффективных систем; защита от дискриминации со стороны алгоритмов; защита конфиденциальности данных; объяснение пользователям работы системы и уведомления о значимых действиях; возможность выбрать человека в качестве альтернативы машине.

В документе отмечается, что он не имеет силы закона, а объединяет принципы, которые нужно использовать при разработке государственных систем, которые затрагивают права человека, а также рекомендуются всем разработчикам автоматизированных систем. 29.10.2022 Синди Гордон, генеральный директор компании Sales Choice, опубликовала в журнале «Форбс» статью с утверждением, что «искусственный интеллект открывает мир бесконтрольных данных, которые подрывают безопасность частной жизни человека. Но если США смогут стать мировым лидером в области управления искусственным интеллектом, это будет способствовать повышению капитализации и получению большей выгоды от технологий ИИ»¹⁶.

Видно, что американские компании и правительство США привержены инновациям и развитию ИИ. В последнее время в США также все

¹⁴ Maintaining American Leadership in Artificial Intelligence: Executive Order 13859 of February 11, 2019. Available at: URL: <https://www.federalregister.gov/documents/2019/02/14/2019-02544/maintaining-american-leadership-in-artificial-intelligence> (дата обращения: 20.07.2023)

¹⁵ Available at: URL: <https://www.wsj.com/articles/white-house-issues-blueprint-for-an-ai-bill-of-rights-11664921544> (дата обращения: 21.07.2023)

¹⁶ Available at: URL: <https://www.forbes.com/sites/cindygordon/2022/10/29/the-usa-blueprint-for-an-ai-rights-advances-usa-ai-governance/?sh=33c52f6950f5> (дата обращения: 20.07.2023)

больше осознают угрозу генеративного ИИ, причем наметилась тенденция к заимствованию европейской модели регулирования: 11.07.2023 в ответ на основные достижения в области технологий генеративного ИИ, а также на важные вопросы, которые эти технологии ставят в таких областях, как интеллектуальная собственность и безопасность человека, Глобальный совет технологической политики Ассоциации вычислительной техники США (АСМ ТРС) выпустил «Принципы разработки, развертывания и использования генеративных технологий искусственного интеллекта». Данный документ совместно подготовлен и принят Комитетом по технологической политике АСМ (USTPC) и Европейским комитетом по технологической политике (Еуроре ТРС). Новые принципы признают, что «растущая мощь систем генеративного ИИ, темпы их эволюции, их широкое применение и потенциал причинения значительного или катастрофического ущерба означают, что необходимо проявлять осторожность при их исследовании, проектировании, разработке, развертывании и использовании. Нынешних механизмов и способов предотвращения такого вреда, вероятно, не будет достаточно»¹⁷. Однако документ является лишь рекомендацией, а не законодательным актом.

2.2. Европейский союз

В Европейском союзе (далее — ЕС) преобладает модель «твердого регулирования», концепция которой в том, чтобы контролировать генеративный ИИ подобно лекарству. Считается, что необходимо создать специальный регулирующий орган, а приложения генеративного ИИ должны пройти тщательное тестирование и получить предварительное одобрение.

В 2020 году Европейская комиссия выпустила «Белую книгу об искусственном интеллекте: европейский подход к совершенству и доверию»¹⁸. По данному документу высоко рискованные приложения, которые могут иметь последствия для прав человека, должны тестироваться и сертифицироваться до поступления на европейский рынок. В апреле 2021 г. Европейская комиссия опубликовала проект Регламента ЕС — «Акта об искусственном интеллекте» (Artificial intelligence act), устанавливающий гармонизированные правила в отношении искусственного ин-

¹⁷ Available at: URL: <https://www.prnewswire.com/news-releases/worlds-largest-association-of-computing-professionals-issues-principles-for-generative-ai-technologies-301873472.html> (дата обращения: 21.07.2023)

¹⁸ Available at: URL: https://ec.europa.eu/info/sites/info/files/commission-whitepaper-artificial-intelligence-feb2020_en.pdf (дата обращения: 21.07.2023)

теллекта¹⁹. Системы высокого риска в проекте «Акта об искусственном интеллекте» получают правила классификации, что, безусловно, стоит признать удачным инструментом правового регулирования. Все приложения ИИ высокого риска должны иметь систему управления рисками, которые помогут определить их качества, исследуемую логику выбора потоков данных, прозрачность и способ обработки направленной пользователем информации.

14.06.2023 депутаты Европарламента одобрили проект «Акта об искусственном интеллекте» 499 голосами против 28 при 93 воздержавшихся. Документ послужил отправной точкой обсуждения окончательной формы закона со странами-членами ЕС. По определению Европейского парламента, принятие законодательства об ИИ направлено на то, чтобы используемые в ЕС системы искусственного интеллекта были «безопасными, прозрачными, отслеживаемыми, недискриминационными и экологически безопасными. Системы ИИ должны контролироваться человеком, а не автоматизированными процессами, чтобы избежать вредных последствий»²⁰. Новые правила устанавливают обязательства провайдеров и пользователей в зависимости от категории риска системы искусственного интеллекта: неприемлемый, высокий, ограниченный, а также генеративный ИИ. Согласно документу, системы генеративного ИИ (такие, как ChatGPT) должны соответствовать требованиям прозрачности (т.е. раскрывать, когда контент был сгенерирован ИИ, в том числе помогать отличать так называемые «дипфейковые» изображения от реальных) и вводить защиту от создания нелегального контента. Кроме того, необходимо сделать общедоступными подробные сводки данных, защищенных авторским правом, которые используются для обучения этих систем. Окончательный вариант закона должен пройти согласование в Еврокомиссии и Евросовете. Планировалось принятие законопроекта до конца 2023 г.²¹

Однако 30.06.2023 руководители крупнейших компаний, включая Siemens и Heineken, направили в Европейскую комиссию, Европарламент письмо против закона о регулировании искусственного интеллекта. Авторы письма особенно волнует генеративный ИИ, поскольку

¹⁹ 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. Available at: URL: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A52021PC0206> (дата обращения: 19.07.2023)

²⁰ Available at: URL: <https://mediacritica.md/ru/голосование-в-европейском-парламент/> (дата обращения: 20.07.2023)

²¹ Available at: URL: <https://ru.euronews.com/my-europe/2023/06/14/ru-eu-artificial-intelligence-vote> (дата обращения: 20.07.2023)

новые правила будут жестко регулировать чат-ботов вне зависимости от цели их использования. Компании, разрабатывающие и внедряющие такие системы, столкнутся с непропорционально высокими затратами на соблюдение требований, говорится в письме²².

2.3. Китайская Народная Республика

Опираясь на основные модели правового регулирования в вышеупомянутых странах и на собственный опыт развития генеративных технологий ИИ, Китай предложил модель «нейтрального регулирования». 11.04.2023 Управление киберпространства КНР издало «Меры управления услугами генеративного искусственного интеллекта (проект для общественного обсуждения)» (далее — «Проект мер»)²³. После запроса и всестороннего рассмотрения общественных мнений 13.07.2023 Управление киберпространства КНР совместно с другими ведомствами опубликовало новые «Временные меры по управлению услугами генеративного искусственного интеллекта» (далее — «Временные меры»)²⁴, вступающие в силу 15 августа. «Временные меры» состоят из пяти разделов: это «общие положения», «развитие технологий и управление ими», «спецификации услуг», «надзор, проверка и юридическая ответственность», а также «дополнительные положения». Что касается регулирования и управления ИИ, то Китай ясно дал понять, что «развитию и безопасности» и также «инновации и управление» должны придаваться одинаковые значения, поэтому автор настоящей статьи называет такую законодательство «нейтральным регулированием».

Во-первых, сфера действия «Временных мер» ограничена сервисами на основе генеративного ИИ, предназначенными для «широкой ответственности», и не включает проекты в области науки, техники и образования с ограниченным доступом. Пункт 1 статьи 2 «Временных мер» гласит: «Настоящие меры применяются к использованию технологий генеративного ИИ для оказания населению Китайской Народной Республики услуг, генерирующих текст, изображения, аудио, видео и другой контент (далее — услуги генеративного ИИ)». По сравнению с «Проектом мер», «Временные меры» дополняют ст. 2 новым пунктом: «Если

²² Available at: URL: <https://incruasia.ru/news/150-kompanij-protiv-zakona-ii/> (дата обращения: 17.07.2023)

²³ Available at: URL: http://www.cac.gov.cn/2023-04/11/c_1682854275475410.htm (дата обращения: 20.06.2023)

²⁴ Временные меры по управлению услугами генеративного искусственного интеллекта. Available at: URL: <https://baijiahao.baidu.com/s?id=1771288400061005032&wfr=spider&for=pc> (дата обращения: 18.07.2023)

отраслевые организации, предприятия, образовательные и научно-исследовательские учреждения, государственные учреждения культуры, соответствующие профессиональные учреждения и т.д. разрабатывают и применяют технологии генеративного ИИ и не предоставляют услуги генеративного ИИ населению на территории КНР, то положения настоящих Мер не применяются».

Такие положения отражают, что законодатель в известной мере исключил из сферы регулирования генеративные сервисы ИИ, которые разрабатываются самостоятельно, применяются внутри компании и не общедоступны. Это значительно снижает нагрузку соблюдения требований на этапе разработки и исследования сервисов, снимает тревоги многих предприятий насчет соблюдения требований при обращении к услугам генеративного ИИ в целях внутреннего применения, например, для повышения эффективности работы, и отражает курс «Временных мер» на поощрение инноваций. Так как генеративный ИИ стал ключевой позицией конкуренции стран в области искусственного интеллекта, речь идет не только о технологическом суверенитете и цифровом суверенитете, но и о будущей промышленной системе и даже о всесторонней мощи страны. Поэтому на современном этапе правительство КНР оставляет технологическим инновациям в области искусственного интеллекта большую свободу действий и пространство для «проб и ошибок».

Во-вторых, классифицированный и иерархический надзор с акцентом на координацию разных органов. Статья 3 «Временных мер» устанавливает, что: «Государство придерживается принципа равноценного значения развития и безопасности, поощрения инноваций и управления в соответствии с законом, принимает меры к развитию генеративного ИИ и осуществляет учитывающий, разумный, классифицированный и иерархический надзор за услугами генеративного ИИ». Пункт 2 статьи 16 также гласит, что «компетентные государственные органы с учетом особенностей технологии генеративного ИИ и применения его услуг в соответствующих отраслях и сферах совершенствуют методы научного регулирования, совместимые с инновациями и развитием, и формулируют правила или руководства по классифицированному и иерархическому надзору».

Как видим, оба приведенных пункта предусматривают, что для генеративного ИИ будет введен «классифицированный и иерархический надзор». Очевидно, что такая идея регулирования опирается на положения проекта «Акта об искусственном интеллекте» ЕС. Кроме того, «Временные меры» затрагивают некоторые ключевые области применения услуг генеративного ИИ, устанавливая, что «если государство имеет другие правила использования услуг генеративного ИИ для осуществления таких видов деятельности, как публикация новостей, производство

фильмов и телепередач, художественное и литературное творчество, эти правила применяются соответствующим образом» (п. 2 ст. 2).

Что касается координации и сотрудничества между регулирующими органами, то во «Временных мерах» также указано, что «департаменты сетей и информации, развития и реформ, образования, науки и техники, промышленности и информационных технологий, общественной безопасности, радио- и телевидения, печати и публикаций должны в соответствии со своими обязанностями усилить управление услугами генеративного ИИ в соответствии с законом» (п. 1 ст. 16). Привлечение различных отраслевых ведомств для надзора также отражает стремление правительства регулировать различные области применения услуг генеративного ИИ. Пока «Временные меры» не содержат ясных и полных указаний, как классифицировать и оценивать услуги или технологии ИИ, и соответствующие правила и стандарты находятся на стадии доработок. В будущем органы по стандартизации смогут сформулировать более подробные правила классификации и оценки генеративного ИИ, а также специальные нормы областей применения или определенных услуг генеративного ИИ с высоким уровнем риска.

В-третьих, «Временные меры» ослабляют требования допуска провайдеров к данной деятельности. Во «Временных мерах» отсутствует изначальное требование «точности и надежности» контента, создаваемого с использованием генеративного ИИ, которое содержит в ст. 4 «Проекта мер», и тем самым снижается обязанность провайдеров предотвращать создание «ложных сведений».

Поскольку модель «черного ящика» ИИ делает невозможным абсолютно действенный контроль над алгоритмами и генерируемыми результатами услуг ИИ, естественно, провайдер не может полностью обеспечить достоверность и точность генерируемого контента. Проверка провайдером подлинности каждой единицы данных и генерируемого контента приведет к большим операционным издержкам. Вместо него во «Временных мерах» предусмотрено положение: «Оказание и использование генеративных услуг ИИ должно соответствовать законам и правилам, уважать общественную мораль и этику и соблюдать следующие положения:

придерживаться основных социалистических ценностей и не генерировать контент, подстрекающий к подрыву государственной власти, свержению социалистической системы, ставящий под угрозу безопасность и интересы страны, наносящий ущерб репутации страны, подстрекающий к расколу государства, подрывающий национальное единство и социальную стабильность, пропагандирующий терроризм, экстремизм, национальную ненависть, этническую дискриминацию, насилие, непристойность и порнографию, ложную и вредную информацию и другой контент, запрещенный законами и правилами;

в процессе разработки алгоритмов, подбора обучающих данных, создания и оптимизации моделей и оказания услуг принимать меры к предотвращению дискриминации по национальному, религиозному, страновому, географическому, половому, возрастному, профессиональному, медицинскому и другим признакам;

соблюдать права интеллектуальной собственности, деловую этику, хранить коммерческую тайну, не использовать алгоритмы, данные, платформы и другие преимущества для осуществления монополии и недобросовестной конкуренции;

уважать законные права и интересы других лиц, не подвергать опасности их физическое и психическое здоровье, не нарушать их права на изображение, репутацию, честь, неприкосновенность частной жизни и личную информацию;

в зависимости от вида сервиса принимать меры к повышению прозрачности услуг генеративного ИИ и улучшения точности и надежности генерируемого контента».

Теперь во «Временных мерах» отсутствует изначальное требование — для обнаруженного в процессе эксплуатации и зафиксированного пользователями сгенерированного контента, не соответствующего требованиям данного метода, помимо таких мер, как контентная фильтрация, в течение трех месяцев должно быть предотвращено его повторное генерирование путем обучения оптимизации модели и другими способами» (ст. 15 «Проекта мер»). В окончательной редакции говорится: «Провайдеры в ходе оказания услуг должны обеспечивать их безопасность, стабильность и непрерывность для поддержания их нормального применения пользователями. Если провайдер обнаруживает незаконный контент, он должен незамедлительно принять меры к его ликвидации (остановка генерации, прекращение передачи, устранение и т.д.), к оптимизации модели и к обучению для исправления ситуации, и сообщить о произошедшем в соответствующие органы. Если провайдер обнаружит, что пользователь использует услугу генеративного ИИ для незаконной деятельности, он должен сделать ему предупреждение, ограничить функции, приостановить или прекратить оказание ему услуг и принять другие меры к ликвидации нарушений в соответствии с законом, в то же время вести соответствующие записи и сообщать в компетентные органы о нарушениях. Провайдеры должны создавать и совершенствовать механизм подачи жалоб и сообщений, создать удобный портал для подачи жалоб и сообщений, публиковать процесс обработки и сроки обратной связи, своевременно принимать и обрабатывать жалобы и сообщения населения и сообщать о результатах» (ст. 13–15 «Временных мер»).

Кроме того, в соответствии с изначальным проектом до оказания услуг для общественности с использованием технологий генеративно-

го ИИ провайдер должен пройти оценку безопасности и регистрацию алгоритмов. Во «Временных мерах» оценка безопасности и регистрация алгоритмов необходима только при оказании услуг, которые обладают «свойствами общественного обсуждения» или «потенциалом мобилизации общества» (ст. 17)²⁵.

В-четвертых, «Временные меры» допускают возможность создания сервисов генеративного ИИ с иностранным участием. Статья 20 гласит: «Если оказание услуг генеративного ИИ на территории, исходящее из-за пределов КНР, не соответствует законам, административным правилам и положениям настоящих мер, Управление киберпространством уведомляет соответствующие организации о необходимости технических мер и других действий». В этом случае компетентные органы вправе принять технические меры, например, блокировку, чтобы закрыть доступ к соответствующим сайтам и приложениям оффшорных сервисов. Если отечественный провайдер внедряет оффшорный генеративный ИИ-сервис в собственный продукт для оказания услуг соотечественникам, он также обязан соблюдать соответствующие положения «Временных мер», в противном случае компетентные органы могут наложить на отечественного провайдера штрафные санкции в соответствии со ст. 21 «Временных мер». В статью 23 «Временных мер» также добавлено: «Если законы и правила предусматривают, что для оказания услуг генеративного ИИ необходимы лицензии, провайдер должен их получить. Иностранные инвестиции в услуги генеративного ИИ должны соответствовать законам и правилам об иностранных инвестициях».

Пока законы и подзаконные акты не устанавливают лицензирования или ограничений доступа иностранных инвестиций к услугам генеративного ИИ как таковых, но если услуги генеративного ИИ применяются в областях, где существует лицензирование, или если они применяются в областях, запрещенных положениями «Специальных мер ввода иностранных инвестиций (негативный список; редакция 2021 года)»²⁶, то они должны соответствовать данным положениям.

С другой стороны, существуют требования к получению таких квалификаций, как лицензия на эксплуатацию сетевой культуры, лицензия на услуги сетевой публикации и лицензия на распространение аудиовизуальных программ в информационной сети, где действуют ограничения или запреты на иностранные инвестиции. Если услуги генеративного

²⁵ Портал системы регистрации алгоритмов. Available at: URL: <https://beian.cac.gov.cn> (дата обращения: 10.07.2023)

²⁶ Available at: URL: https://www.gov.cn/zhengce/2022-11/28/content_5713317.htm (дата обращения: 10.07.2023)

ИИ применяются в таких областях, то на них могут распространяться определенные ограничения. Поэтому иностранцам, инвестирующим в эти области, следует всесторонне изучить отрасли и требования к входным квалификационным характеристикам инвестиционных проектов. Поскольку зарубежная индустрия генеративного ИИ занимает лидирующее технологическое положение, технологическое развитие отечественных производителей все еще невысоко, а приведенные выше положения об иностранных инвестициях в области генеративного ИИ остаются неясными, несовершенство мер может увеличить риск прекращения технической поддержки из-за рубежа или принудительной передачи технологий.

3. Основные тенденции совершенствования/ правового регулирования сервисов генеративного ИИ в Китае

3.1. Развитие инноваций в области технологии искусственного интеллекта

Обращаясь к анализу разных моделей правового регулирования сервисов генеративного ИИ в современном мире, необходимо отметить, что в КНР основным ориентиром преобразований будет служить инновация в области технологии искусственного интеллекта [Гу Н., 2023: 83]. 13.03.2023 Пекинское муниципальное бюро экономики и информатизации выпустило «Белую книгу о развитии индустрии искусственного интеллекта в Пекине в 2022 году», обозначившую цели развития генеративного ИИ: «Всестороннее развитие индустрии искусственного интеллекта, поддержание предприятий в создании крупных ИИ-моделей, сравнимых с ChatGPT, укрепление инфраструктуры, ускорение предоставления базовых данных и создание платформ для сотрудничества фирм, университетов и разработчиков ПО с открытым исходным кодом»²⁷. Единственным способом повышения качества и точности информации остается совершенствование алгоритмов. Как отмечают специалисты, «точность фильтрации контента, которая в основном состоит из фильтрации по ключевым словам, в значительной степени зависит от совершенствования алгоритмической технологии» [Сунь Я., Чжоу С., 2011: 45].

Поскольку генеративный ИИ по своей сути является специализированным и сложным, его управление должно включать участие множес-

²⁷ Available at: URL: https://www.beijing.gov.cn/ywdt/gzdt/202302/t20230214_2916514.html?eqid=8e6b9d6f000dcb61000000026479c24f (дата обращения: 04.07.2023)

тва субъектов, включая правительства, предприятия и граждан [Ван Я., Янь Х., 2023: 21]. Жесткие законодательные и нормативные требования к предприятиям, находящимся на ранних стадиях технологического развития, могут сдерживать и тормозить инновации в области ИИ, а также ограничивать способность всех слоев общества быстро реагировать на будущие технологические прорывы и достижения. Требование к технологическим компаниям разрабатывать технологии с соблюдением этических норм может оказаться более действенным, чем наложение на них многомиллионных штрафов. В рамках недавнего судебного процесса против компании Meta Министерство юстиции США обязало Meta к концу 2022 года разработать инструменты устранения алгоритмической дискриминации в персонализированной рекламе. В январе 2023 года компания ввела в эксплуатацию систему Variance Reduction, снижающую риск гендерного и расового неравенства в рекламе с помощью методов дифференцированной конфиденциальности²⁸.

Кроме того, нельзя недооценивать полезности участия в управлении генеративным ИИ общественности и технологических сообществ. В начале февраля 1923 года Минцифры России запустило проект поиска уязвимостей на Госуслугах. Три месяца более 8 тыс. участников проверяли защищенность портала, их работа помогла улучшить систему безопасности Госуслуг²⁹. Поэтому следует мобилизовать силы всего общества, чтобы добиться совместного повышения инновационного и регулятивного потенциала.

4. Заимствование опыта российского законодательства в области защиты данных

Внесенные в 2020 г. в Конституцию России поправки относят к сфере федеральной компетенции «обеспечение безопасности личности, общества и государства при применении информационных технологий, обороте цифровых данных» (п. «м» ст. 71)³⁰. Прежняя редакция Конституции к федеральной компетенции относила лишь «федеральную информацию и связь». Президент России также подписал Указ, закрепляю-

²⁸ Justice Department and Meta Platforms Inc. reach key agreement as they implement groundbreaking resolution to address discriminatory delivery of housing advertisements, Jan. 9, 2023. Available at: URL: <https://www.justice.gov/opa/pr/justice-department-and-meta-platforms-inc-reach-key-agreement-they-implement-groundbreaking> (дата обращения: 10.07.2023)

²⁹ Available at: URL: <https://digital.gov.ru/ru/events/44375/> (дата обращения: 09.07.2023)

³⁰ Конституция Российской Федерации (принята на всенародном голосовании 12.12.1993) (с изменениями, одобренными в ходе общероссийского голосования 1.07.2020) // СПС КонсультантПлюс.

щий, что «с 31.03.2022 заказчикам (кроме организаций с муниципальным участием) запрещено закупать иностранное программное обеспечение для использования на значимых объектах критической информационной инфраструктуры, а также услуги по использованию такого ПО без согласования с уполномоченным органом. С 1.01.2025 органам власти и заказчикам запрещается использовать иностранное ПО на значимых объектах критической информационной инфраструктуры»³¹.

Помимо установления ценности гарантий информационной безопасности, ориентированных на человека, 30.04.2023 Правительство России также утвердило новую «Концепцию информационной безопасности детей» в России. Концепция предусматривает, в частности, «введение в российских школах уроков информационной безопасности и цифровой грамотности; просветительские мероприятия для родителей, учителей, работников детских и юношеских библиотек; расширение спектра возможностей услуг родительского контроля на стационарных и мобильных устройствах, которыми пользуется ребенок»³².

В настоящее время законодательство КНР о безопасности данных содержит много принципиальных положений и мало конкретных мер борьбы с угрозами информационной безопасности, особенно в области генеративного ИИ. Если иностранные правительства или связанные с ними организации используют генеративные модели искусственного интеллекта для передачи информации, нарушающей законы и правила Китая, то китайское правительство может применить не только технические меры блокирования передачи, но и более широкие правовые контрмеры для лучшей защиты цифрового суверенитета, безопасности данных, как в России.

5. Укрепление международного сотрудничества в цифровой сфере

Поскольку рабочий механизм сервисов генеративного ИИ основан на потоке данных, а правила раскрытия данных и операционной совместимости в настоящее время различаются в разных странах, необходимо разработать единые правила раскрытия данных для международного сотрудничества. Международные организации должны способствовать разработке соглашений о регулировании генеративного ИИ. Только сов-

³¹ Указ Президента РФ от 30.03.2022 № 166 “О мерах по обеспечению технологической независимости и безопасности критической информационной инфраструктуры Российской Федерации” // СЗ РФ. 2022. N 14. Ст. 2242.

³² Available at: URL: <https://www.garant.ru/products/ipo/prime/doc/406740607/?ysclid=lig0cuz7rt451127258> (дата обращения: 04.07.2023)

местная разработка стандартов и процедур управления рисками генеративного ИИ и продвижение согласованного на международном уровне законодательства позволит справиться с будущими межгосударственными спорами, вызванными генеративным ИИ, и поэтапно построить более широкое и всеобъемлющее международное сотрудничество. Сейчас Китай активно участвует в консультациях и обсуждениях по цифровым вопросам в ООН, ВТО, G20, АТЭС, БРИКС, ШОС и других платформах, продвигает «Рамочное соглашение партнерства цифровой экономики БРИКС», «Инициативу взаимодействия в обеспечении безопасности цифровых данных между Китаем и странами Центральной Азии» и другие инициативы.

Кроме того, Китай заявил о создании международной организации для деятельного участия в управлении Интернетом³³. На сегодняшний день ее членами являются более 20 стран на шести континентах, включая более 100 учреждений, организаций, предприятий и частных лиц. В будущем Китай продолжит углублять двустороннее и многостороннее сотрудничество в области международной информационной безопасности, а также постепенно распространит опыт сотрудничества и его уже действующие механизмы на взаимодействие в области информационной безопасности: с одной стороны, необходимо укреплять доверие, обмен и правовое сотрудничество в области информационной безопасности, расширять обмен разведанными между правоохранительными органами и службами безопасности. С другой стороны, есть необходимость активно использовать ШОС и другие платформы сотрудничества в упрочении безопасности данных, проводить совместные исследования и модернизировать инфраструктуру безопасности данных, наращивать совместные учения в области наступательных и оборонительных возможностей киберпространства, укреплять способность всех сторон действовать в чрезвычайных ситуациях. Сотрудничество правительств России и Китая в совместном управлении рисками безопасности данных генеративного ИИ предстоит продолжать.

Заключение

Способность понимать человеческий язык и запоминать большое количество фактов, полученных в процессе обучения, подкрепляется огромными объемами данных и мощными вычислительными мощностями, что позволяет технологиям генеративного ИИ не только понимать

³³ Available at: URL: <https://www.iksmedia.ru/news/5895608-Kitaj-zayavil-o-sozdanii-mezhdunar.html?ysclid=ljw5xehbk992655583> (дата обращения: 01.07.2023)

естественный человеческий язык, но и генерировать высококачественный контент на основе «запомненных» знаний. Это значит, что потенциальные риски, создаваемые технологиями генеративного ИИ для человеческого общества, актуальнее нынешних. Развитие науки и техники — это обоюдоострый меч, и хотя мы наслаждаемся удобством ИИ, мы также должны опасаться его влияния на традиционную правовую систему и твердо управлять им.

Конечно, правовые проблемы, которые мы теперь наблюдаем, могут быть лишь верхушкой айсберга, но создание соответствующей правовой системы — это не ограничение или сопротивление генеративному ИИ, а минимизация угрозы, которой он может стать человечеству. В настоящее время В приложении к «Плану законодательной работы Госсовета КНР на 2023 год», выпущенном 31.05.2023, указано что проект Закона «Об искусственном интеллекте» будет вынесен на рассмотрение Постоянного комитета Всекитайского собрания народных представителей КНР. Кроме того, Национальный технический комитет стандартизации информационной безопасности КНР недавно активизировал разработку двух стандартов генеративного ИИ, включая «Спецификацию безопасности данных предварительного обучения и оптимизационного обучения генеративного искусственного интеллекта» и «Спецификацию безопасности для искусственной маркировки генеративного искусственного интеллекта». Как сочетать верховенство закона и инновации и способствовать здоровому развитию и стандартизированному применению генеративного ИИ — это значимый вопрос, стоящий перед законодателем.



Список источников

1. 王洋、闫海：《生成式人工智能的风险迭代与规制革新——以ChatGPT为例》，载《理论月刊》2023年第6期，第14–24页。[Ван Я., Янь Х. Итерация рисков и регуляторные инновации генеративного искусственного интеллекта — на примере ChatGPT // Ежемесячная теория. 2023. N 6. С. 14–24].
2. 顾男飞：《生成式人工智能的智能涌现、风险规制与产业协调》，载《荆楚法学》2023年第3期，第70–83页。[Гу Н. Эмерджентная способность, регулирование рисков и промышленная координация генеративного искусственного интеллекта // Цзиньчжу правоведение. 2023. N 3. С. 70–83].
3. 邓建鹏、朱悸成：《ChatGPT模型的法律风险及应对之策》，载《新疆师范大学学报》《哲学社会科学版》2023年第5期，第91–101页。[Дэн Ц., Чжу И. Правовые риски модели ChatGPT и стратегии их преодоления // Вестник Синьцзянского педагогического университета (Философия и социальные науки). 2023. N 5. С. 91–101].
4. 刘艳红：《生成式人工智能的三大安全风险及法律规制——以ChatGPT为例》，载《东方法学》2023年第4期，第30–43页。[Лю Я. Три основных риска безопаснос-

ти генеративного искусственного интеллекта и правовое регулирование — на примере ChatGPT // Восточное право. 2023. N 4. С. 30–43.]

5. 孙艳、周学广：《内容过滤技术研究进展》，载《信息安全与通信保密》2011年第9期，第45–49页。[Сунь Я., Чжоу С. Исследование прогресса в технологии фильтрации контента // Информационная безопасность и тайна связи. 2011. N 9. С. 45–49.]

6. 曹建峰：《迈向可信AI：ChatGPT 类生成式人工智能的治理挑战及应对》，载《上海政法学院学报》2023年第4期，第28–42页。[Цао Ц. На пути к надежному ИИ: проблемы управления и ответы на генеративный искусственный интеллект, подобный ChatGPT // Вестник Шанхайского института политики и права. 2023. N 4. С. 28–42.]

7. 曹树金、曹茹焱：《从ChatGPT看生成式AI对情报学研究与实践的影响》，载《现代情报》2023年第4期，第3–10页。[Цао Ш., Цао Р. Влияние генеративного ИИ на исследования и практику в области науки о разведке на примере ChatGPT // Современная наука о разведке. 2023. N 4. С. 3–10.]

8. 蔡士林、杨磊：《ChatGPT智能机器人应用的风险与协同治理研究》，载《情报理论与实践》，2023年第5期。第14–22页。[Цай Ш., Ян Л. Исследование рисков и совместного управления приложением интеллектуального робота ChatGPT // Теория и практика разведки. 2023. N 5. С. 14–22.]

9. 靳思远：《全球数据治理的DEPA路径和中国的选择》，载《财经法学》2022年第6期，第96–110页。[Цзинь С. Путь DEPA глобального управления данными и выбор Китая // Финансово-экономическое право. 2022. N 6. С. 96–110.]

10. 阙天舒、吕俊延：《智能时代下技术革新与政府治理的范式变革》，载《中国行政管理》，2021年第2期，第21–30页。[Цюэ Т., Лу Ц. Технологические инновации и изменение парадигмы государственного управления в эпоху интеллекта // Китайское административное управление. 2021. N 2. С. 21–30.]

11. 支振锋：《生成式人工智能大模型的信息内容治理》，载《政法论坛》2023年第4期，第34–48页。[Чжи Ч. Генеративные большие модели искусственного интеллекта для управления информационным контентом // Форум политики и права. 2023. N 4. С. 34–48.]

12. 陈兵：《生成式人工智能可信发展的法治基础》，载《上海政法学院学报》2023年第4期，第13–27页。[Чэнь Б. Основа верховенства права для достоверного развития генеративного искусственного интеллекта // Вестник Шанхайского института политики и права. 2023. N 4. С. 13–27.]

13. 商建刚：《生成式人工智能风险治理元规则研究》，载《东方法学》2023年第3期，第4–17页。[Шань Ц. Исследование правил для управления рисками генеративного ИИ // Восточное право. 2023. N 3. С. 4–17.]

14. 袁曾：《生成式人工智能的责任能力研究》，载《东方法学》2023年第3期，第18–33页。[Юань Ц. Исследование способности ответственности генеративного искусственного интеллекта // Восточное право. 2023. N 3. С. 18–33.]

15. 於兴中、郑戈、丁晓东：《生成性人工智能与法律：以ChatGPT为例》，载《中国法律评论》2023年第2期，第1–20页。[Юй С., Чжэн Г., Дин С. Генеративный искусственный интеллект и право: на примере ChatGPT // Китайский юридический обзор. 2023. N 2. С. 1–20.]

16. Wei J., Tay Y. et al. Emergent abilities of large language models. Transactions on Machine Learning Research, 2022, no. 8, pp. 1–2.



References

1. Cai S., Yang L. (2023) Study in the risk and collaborative governance of ChatGPT intelligent robot application. *Qing bao li lun yu shi jian*=Intelligence Theory and Practice, no 5, pp. 14–22 (in Chinese)
2. Cao J. (2023) Toward trustworthy AI: governance challenges and responses to ChatGPT-like generative artificial intelligence. *Shang hai zheng fa xue yuan xue bao*=Bulletin of Shanghai Institute of Politics and Law, no. 4, pp. 28–42 (in Chinese)
3. Cao S., Cao R. (2023) Impact of generative AI on exploration science research and practice using ChatGPT as an example. *Xian dai qing bao*=Contemporary Exploration Science, no. 4, pp. 3–10 (in Chinese)
4. Chen B. (2023) Rule of law basis for the credible development of generative artificial intelligence. *Shang hai zheng fa xue yuan xue bao*=Bulletin of Shanghai Institute of Politics and Law, no. 4, pp. 13–27 (in Chinese)
5. Deng J., Zhu Y. (2023) Legal risks of ChatGPT model and strategies to overcome them. *Xin jiang shi fan da xue xue bao (zhe xue she ke ban)*=Bulletin of Xinjiang Pedagogical University (Philosophy and Social Science Series), no. 5, pp. 91–101 (in Chinese)
6. Gu N. (2023) Emergent ability, risk regulation and industrial coordination of generative artificial intelligence. *Jin chu fa xue*=Jingchu Law Review, no. 3, pp. 70–83 (in Chinese)
7. Jin S. (2022) The DEPA path of global data governance and China's choice. *Cai jing fa xue*=Finance and Economics Law, no. 6, pp. 96–110 (in Chinese)
8. Liu Y. (2023) The three major security risks of generative artificial intelligence and legal regulation — Taking ChatGPT as an Example. *Dong fang fa xue*=Eastern Law, no. 4, pp. 30–43 (in Chinese)
9. Que T., Lu C. (2021) Technological innovation and the paradigm shift of public administration in the intelligence era. *Zhong guo xing zheng guan li*=Chinese Administrative Management, no. 2, pp. 21–30 (in Chinese)
10. Shang J. (2023) Study in rules for generative AI risk governance. *Dong fang fa xue*=Eastern Law, no. 3, pp. 4–17 (in Chinese)
11. Sun Y., Zhou X. (2011) Study in progress of content filtering technology. *Xin xi an quan yu tong xin bao mi*=Information Security and Communication Secrecy, no. 9, pp. 45–49 (in Chinese)
12. Wang Y., Yan H. (2023) Risk iteration and regulatory innovation of generative artificial intelligence — a case study of ChatGPT. *Li lun yue kan*=Monthly Theoretical Journal, no. 6, pp. 14–24 (in Chinese)
13. Wei J., Tay Y. et al. (2022) Emergent abilities of large language models. *Transactions on Machine Learning Research*, no. 8, pp. 1–2.
14. Yu X., Zheng G., Ding X. (2023) Generative Artificial Intelligence and the Law: Taking ChatGPT as an Example. *Zhong guo fa lv ping lun*=China Law Review, no. 2, pp. 1–20 (in Chinese)
15. Yuan Z. (2023) Study in the responsibility ability of generative artificial intelligence. *Dong fang fa xue*=Eastern Law, no. 3, pp. 18–33 (in Chinese)
16. Zhi Z. (2023) Generative artificial intelligence large models for information content governance. *Zheng fa lun tan*=Politics and Law Forum, no. 4, pp. 34–48 (in Chinese)

Информация об авторе:

Ли Яо — аспирантка.

Information about the author:

Li Yao — Postgraduate Student.

Статья поступила в редакцию 10.06.2023; одобрена после рецензирования 14.07.2023; принята к публикации 14.07.2023.

The article was submitted to editorial office 10.07.2023; approved after reviewing 14.07.2023; accepted for publication 14.07.2023.