

Дата поступления рукописи в редакцию: 14.06.2025 г.
Дата принятия рукописи в печать: 30.06.2025 г.

DOI: 10.24412/2076-1503-2025-6-438-442

КОГТЕВА Анна Николаевна,
кандидат экономических наук,
доцент, доцент кафедры
информационной безопасности
Финансового Университета
при Правительстве РФ,
e-mail: annelya1@yandex.ru

ВОЛЧЕНКО Анастасия Владимировна,
кандидат юридических наук, старший преподаватель
кафедры уголовного процесса Белгородского юридического
института МВД России имени И.Д. Путилина, майор полиции,
e-mail: espy207@yandex.ru

КОНФИДЕНЦИАЛЬНОСТЬ ПАЦИЕНТОВ КАК КЛЮЧЕВАЯ ПРОБЛЕМА ПРИ ИСПОЛЬЗОВАНИИ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА В МЕДИЦИНСКОЙ СФЕРЕ

Аннотация. В статье рассматривается проблема конфиденциальности пациентов при внедрении искусственного интеллекта в здравоохранение. Цель работы – проанализировать международный опыт правового регулирования применения искусственного интеллекта в медицине с акцентом на защиту персональных данных. Рассматриваются различные подходы к решению этой проблемы: США опираются на закон HIPAA и рекомендации FDA; Евросоюз ввёл всеобъемлющий Акт об искусственном интеллекте (2025 г.) в дополнение к строгим нормам GDPR; Китай делает упор на локальную обработку медицинских данных и собственные стандарты. Отдельно отмечены глобальные инициативы ВОЗ и ОЭСР, задающие общие принципы этичного и безопасного использования искусственного интеллекта. Приведены показательные примеры: сотрудничество британской NHS с компанией DeepMind (нарушение конфиденциальности пациентов привело к реформам) и международный проект AI for COVID с федеративным обучением без передачи данных. В заключение подчеркивается, что заимствование успешных зарубежных практик способствует формированию безопасных и этически устойчивых систем, адаптированных к национальному контексту. Указывается на необходимость интеграции требований безопасности и конфиденциальности на ранних этапах разработки проектов, а также обязательного участия экспертов по кибербезопасности в процессе внедрения. Отдельное внимание уделяется соблюдению международных и национальных норм защиты персональных данных и этических стандартов. Обоснована целесообразность использования регуляторных песочниц для апробации инновационных решений, а также важность международного сотрудничества для стандартизации и обмена опытом. Итоговый акцент сделан на балансе между технологическим развитием и защитой прав пациентов, что обеспечивает устойчивость и эффективность медицинской инфраструктуры при интеграции технологий искусственного интеллекта. Предложены рекомендации для безопасного и этичного внедрения искусственного интеллекта в медицину.

Ключевые слова: искусственный интеллект; здравоохранение; медицинская сфера; информационная безопасность; конфиденциальность; конфиденциальность пациентов; защита персональных данных; этические принципы.

KOGTEVA Anna Nikolaevna,
Candidate of Economic Sciences, Associate Professor,
Associate Professor of the Department
of Information Security
University of Finance
under the Government of the Russian Federation,

VOLCHENKO Anastasia Vladimirovna,
Candidate of Law, Senior Lecturer
at the Department of Criminal Procedure
of the I.D. Putilin Belgorod Law Institute of the Ministry
of Internal Affairs of Russia,
Police Major

PATIENT CONFIDENTIALITY AS A KEY ISSUE IN THE USE OF ARTIFICIAL INTELLIGENCE IN THE MEDICAL FIELD

Annotation. The article discusses the problem of patient confidentiality in the implementation of artificial intelligence in healthcare. The purpose of the work is to analyze the international experience of legal regulation of the use of artificial intelligence in medicine with an emphasis on the protection of personal data. Various approaches to solving this problem are being considered: the United States relies on HIPAA law and FDA recommendations; the European Union has introduced a comprehensive Artificial Intelligence Act (2025) in addition to strict GDPR standards; China focuses on local medical data processing and its own standards. The global initiatives of the WHO and the OECD, which set general principles for the ethical and safe use of artificial intelligence, were highlighted separately. Illustrative examples are given: the cooperation of the British NHS with DeepMind (violation of patient confidentiality led to reforms) and the international AI for COVID project with federated training without data transfer. In conclusion, it is emphasized that borrowing successful foreign practices contributes to the formation of safe and ethically sustainable systems adapted to the national context. It is pointed out that it is necessary to integrate security and confidentiality requirements at the early stages of project development, as well as the mandatory participation of cybersecurity experts in the implementation process. Special attention is paid to compliance with international and national standards of personal data protection and ethical standards. The expediency of using regulatory sandboxes for testing innovative solutions is substantiated, as well as the importance of international cooperation for standardization and exchange of experience. The final focus is on the balance between technological development and the protection of patients' rights, which ensures the sustainability and efficiency of the medical infrastructure while integrating artificial intelligence technologies. Recommendations for the safe and ethical implementation of artificial intelligence in medicine are proposed.

Key words: artificial intelligence; healthcare; medical field; information security; confidentiality; patient confidentiality; personal data protection; ethical principles.

Введение. Широкое внедрение искусственного интеллекта в медицину обострило вопрос соблюдения конфиденциальности пациентов. Алгоритмы искусственного интеллекта обучаются на больших массивах медицинских данных, а существующие механизмы защиты информации не успевают за стремительным технологическим прогрессом. Это создаёт повышенный риск утечек конфиденциальных сведений и нарушения врачебной тайны.

Проблема исключительно актуальна: резонансные случаи несанкционированной передачи данных (например, без согласия пациентов) подтверждают реальность угрозы. Если не принять надлежащие меры безопасности и регулирования, под ударом окажутся доверие к медицинским инновациям и готовность общества принимать искусственный интеллект в здравоохранении.

Основная часть исследования. Одним из ключевых вопросов остается соблюдение требований конфиденциальности пациентов при

использовании искусственного интеллекта. Международный опыт показывает разнообразие подходов к правовому регулированию этой сферы.

Так, в США защита медицинской информации традиционно обеспечивается законом HIPAA, устанавливающим строгие стандарты обращения с персональными медицинскими данными. Однако появление больших данных и искусственного интеллекта порождает новые сценарии использования информации, которые не полностью обеспечиваются прежними нормами. В последние годы американские регуляторы неоднократно выпускали рекомендации по обеспечению прозрачности алгоритмов и недопущению дискриминации, а Управление по санитарному надзору за качеством пищевых продуктов и медикаментов (FDA) разработало подходы к регулированию программного обеспечения как медицинского изделия (Software as a Medical Device, SaMD), куда включены и алгоритмы искусственного интеллекта. FDA требует от разработчиков доказа-

тельств безопасности и эффективности алгоритмов, включая описание процессов обновления моделей и защиты от несанкционированного вмешательства в их работу.

К настоящему времени запущены пилотные программы по сертификации адаптивных алгоритмов (так называемых *learning algorithms*) с целью контролировать изменения модели после внедрения. Тем не менее, на федеральном уровне в США пока нет единого закона, специально посвященного искусственному интеллекту, поэтому система безопасности основывается во многом на отраслевых стандартах и добровольных инициативах, а также на контроле со стороны профессиональных сообществ (например, Радиологическое сообщество США ведет реестр алгоритмов и публикует рекомендации по их безопасному использованию).

Европейский союз пошел по пути формирования всеобъемлющей нормативной базы в исследуемой сфере. В 2021-2022 гг. Еврокомиссией был представлен проект Акта об искусственном интеллекте (AI Act) – первого в мире масштабного закона, комплексно регулирующего использование искусственного интеллекта во всех отраслях, включая медицину. Обозначенный акт вступил в силу в феврале 2025 г. и ввел категоризацию систем искусственного интеллекта по уровням риска. Соответствующие медицинские системы отнесены к категории высокого риска, что означает повышенные требования к ним. Разработчики и эксплуатанты таких систем обязаны обеспечивать кибербезопасность и устойчивость к внешним угрозам на протяжении всего жизненного цикла изделия. Кроме того, AI Act содержит положения о необходимости внедрения механизмов надлежащего надзора, прозрачности, точности, надежности и кибербезопасности для систем, используемых в сфере здоровья.

Кроме того, вес уже действует строгий Общий регламент по защите данных (GDPR), распространяющийся на любые персональные данные, включая медицинские. Это означает, что проекты искусственного интеллекта в медицине в странах ЕС должны соответствовать принципам минимизации собираемых данных, целевого использования информации и получения информированного согласия пациентов [1]. Например, при развертывании в клинике нового сервиса врач обязан разъяснить пациенту, какие данные будут обработаны алгоритмом и с какой целью.

Европейское агентство по кибербезопасности (ENISA) также выпустило ряд рекомендаций, касающихся защиты медицинских устройств с модулями искусственного интеллекта, учитываю-

щих весь цикл их эксплуатации – от требований к безопасной разработке до пострыночного мониторинга уязвимостей.

Приведенное свидетельствует о переходе ЕС от «мягкого» права (этических принципов) к жесткому законодательному регулированию в области искусственного интеллекта, благодаря чему устанавливаются единые для всех стран-членов правила безопасного использования алгоритмов.

Китай за последние пять лет совершил значительный рывок в нормативном обеспечении искусственного интеллекта в медицине. В 2022 г. Государственное управление по надзору за медицинскими изделиями (NMPA, китайский регулятор медицинских изделий) опубликовало руководящие принципы по оценке и регистрации медицинских изделий на основе машинного обучения. Эти рекомендации требуют от производителей предоставлять подробные данные об используемых наборах для обучения, алгоритмических моделях и мерах по обеспечению кибербезопасности и конфиденциальности.

Китайские правила делают упор на локализации данных – медицинские данные граждан должны обрабатываться внутри страны, а их передача за рубеж строго контролируется. Это связано с общегосударственной политикой киберсуверенитета и стремлением минимизировать внешние риски. Одновременно Китай участвует в международных инициативах по гармонизации подходов: так, NMPA стремится учитывать международный опыт и содействует разработке единых стандартов безопасности искусственного интеллекта вместе с профильными организациями других стран.

В целом, китайский подход сочетает жесткое государственное регулирование и стандартизацию с крупными инвестициями в собственные разработки искусственного интеллекта, включая системы киберзащиты.

Помимо национальных усилий, заметную роль играют международные организации и мультинациональные инициативы. В 2018 г. ВОЗ совместно с Международным союзом электросвязи инициировали Группу по вопросам искусственного интеллекта для здоровья (FG-AI4H), а в 2023 г. на ее основе запущена Глобальная инициатива по искусственному интеллекту в здравоохранении (Global Initiative on AI for Health) [2]. Целью этих программ является выработка общих стандартов и рекомендаций по безопасному и эффективному внедрению искусственного интеллекта, включая этические принципы, вопросы защиты данных и качество алгоритмов.

Организация экономического сотрудничества и развития (ОЭСР) еще в 2019 г. приняла Принципы искусственного интеллекта, поддержанные впоследствии странами G20, которые провозглашают необходимость надежности и безопасности соответствующих систем. В развитие этих принципов в 2023 г. ОЭСР учредила Глобальный форум по искусственному интеллекту, где одной из ключевых тем стала кибербезопасность алгоритмов в таких чувствительных областях, как здравоохранение [3, р. 5].

Также следует упомянуть выпущенные Всемирной организацией здравоохранения в 2021 г. Руководящие принципы по этике и управлению искусственным интеллектом для здравоохранения, где уделено внимание защите персональных данных пациентов и предотвращению дискриминации при использовании соответствующих алгоритмов. Эти документы задают ориентиры для стран, разрабатывающих собственные политики в данной сфере.

Для иллюстрации интеграции искусственного интеллекта и мер безопасности представляется целесообразным рассмотреть несколько показательных примеров.

В Великобритании внимание общественности привлек проект сотрудничества Национальной службы здравоохранения (NHS) с компанией DeepMind (подразделение Google) по анализу медицинских данных для диагностики почечной недостаточности. Проект продемонстрировал высокую эффективность алгоритма, но столкнулся с критикой со стороны регуляторов за нарушение правил конфиденциальности: было установлено, что одна из больниц передала компании данные ~1,6 млн пациентов без надлежащего информирования и согласия. В результате в 2017 г. Управление комиссара по информации (ICO) признало такой перенос данных незаконным и обязало усилить меры защиты и прозрачности в дальнейшем [4].

Этот случай стал поворотным моментом: NHS пересмотрела свои подходы к партнерству с технологическими компаниями, введя более строгие требования к обезличиванию данных и информированию пациентов. С другой стороны, он подтолкнул к разработке четких национальных рекомендаций по этике искусственного интеллекта в здравоохранении в Великобритании, которые были опубликованы в 2019-2020 гг. и легли в основу текущей политики безопасности NHS в соответствующих проектах.

Другой пример – международный проект AI for COVID, в рамках которого в 2020-2021 гг. объединились радиологические службы нескольких

десятков больниц из разных стран (США, Италия, Китай, Россия и др.) для создания алгоритма, распознающего COVID-19 по КТ легких. В этом проекте впервые в столь широком масштабе была применена федеративная модель обучения: сырые данные пациентов не покидали пределы каждой страны, обучающие процессы происходили локально, а агрегировались только обновления модели.

Проект был успешно реализован: полученный алгоритм показал высокую точность диагностики, при этом организаторы заявили об отсутствии утечек данных благодаря изначально заложенному принципу безопасности и приватности. Данный кейс продемонстрировал жизнеспособность подхода совместного развития искусственного интеллекта без нарушения суверенитета данных, что особенно важно в международном контексте, где законодательство о конфиденциальности может различаться.

Обобщение международного опыта показывает, что для успешного и безопасного применения искусственного интеллекта в медицине необходим баланс инноваций и контроля. С одной стороны, регулятивное давление безусловно нужно – оно обеспечивает минимальные стандарты безопасности, формирует доверие пользователей и пациентов. С другой стороны, чрезмерно жесткие ограничения могут тормозить внедрение перспективных технологий. Страны ищут компромисс: например, в ряде юрисдикций вводится понятие «песочницы искусственного интеллекта» (AI sandbox) для здравоохранения – это специальный режим, в котором инновационные продукты могут ограниченно применяться на практике (например, в пилотных больницах или при наблюдении регулятора) до полного получения сертификатов. Такой режим апробирован в Сингапуре, Японии и некоторых других странах, позволяя собрать информацию о реальных рисках и доработать меры безопасности перед масштабированием решения на всю систему здравоохранения.

В то же время международное сотрудничество в области кибербезопасности медицинского искусственного интеллекта расширяется: создаются платформы обмена информацией об инцидентах, уязвимостях и лучших практиках (как, например, рабочая группа по кибербезопасности здравоохранения при Глобальном альянсе по безопасности пациентов). Подобный обмен опытом между странами и организациями помогает быстрее реагировать на новые угрозы и вырабатывать единые стандарты.

Заключение. Практическая ценность изучения международного опыта заключается в воз-

возможности заимствования наиболее успешных решений. Странам, только начинающим масштабное внедрение искусственного интеллекта в медицине, важно учитывать уроки, полученные их предшественниками: необходимость закладывать требования безопасности на этапе планирования проектов, вовлекать экспертов по кибербезопасности в разработку и внедрение систем искусственного интеллекта, обеспечивать соответствие алгоритмов действующим нормам защиты данных и этическим стандартам. Грамотное сочетание технологических инноваций и мер безопасности позволит реализовать огромный потенциал обозначенных систем в здравоохранении без ущерба для прав пациентов и устойчивости медицинской инфраструктуры.

Основные рекомендации по дальнейшему расширению использования искусственного интеллекта в медицине:

- закладывать требования безопасности и конфиденциальности данных на этапе планирования проектов с искусственным интеллектом;
- вовлекать специалистов по кибербезопасности на всех стадиях разработки и внедрения алгоритмов;
- соблюдать действующие нормы защиты персональных данных (например, GDPR, HIPAA) и этические стандарты при разработке и использовании медицинских систем искусственного интеллекта;
- обеспечивать прозрачность работы алгоритмов и недопущение дискриминации при принятии ими решений;
- использовать специальный ограниченный режим внедрения («регуляторная песочница») для тестирования инновационных продуктов искусственного интеллекта в медицине перед их масштабным запуском;

– расширять международное сотрудничество для обмена опытом, оперативного реагирования на инциденты и выработки единых стандартов безопасности;

– поддерживать баланс между технологическими инновациями и надлежащим контролем, чтобы искусственный интеллект приносил пользу без ущерба для прав пациентов и стабильности системы здравоохранения.

Список литературы:

- [1] Palaniappan K.; Lin E.Y.T.; Vogel S. Global Regulatory Frameworks for the Use of Artificial Intelligence (AI) in the Healthcare Services Sector. // Healthcare. – 2024. – №12 (562).
- [2] WHO: Global Initiative on AI for Health // World Health Organization. – 2025. – Mar 25.
- [3] Krunal M.G. The Role of AI in Shaping Global Healthcare Cybersecurity Policies // International Journal For Multidisciplinary Research. – 2024. – №6(5). – 14 p.
- [4] Khan T. Data Breaches and Liability in the Age of AI: Who's responsible? // Barrister Group. – 2024. – 29 October.

Spisok literatury:

- [1] Palaniappan K.; Lin E.Y.T.; Vogel S. Global Regulatory Frameworks for the Use of Artificial Intelligence (AI) in the Healthcare Services Sector. // Healthcare. – 2024. – №12 (562).
- [2] WHO: Global Initiative on AI for Health // World Health Organization. – 2025. – Mar 25.
- [3] Krunal M.G. The Role of AI in Shaping Global Healthcare Cybersecurity Policies // International Journal For Multidisciplinary Research. – 2024. – №6(5). – 14 p.
- [4] Khan T. Data Breaches and Liability in the Age of AI: Who's responsible? // Barrister Group. – 2024. – 29 October.



ЮРКОПАНИ
www.law-books.ru

Юридическое издательство «ЮРКОПАНИ»
издает научные журналы:

- Научно-правовой журнал «Образование и право», рекомендованный ВАК Министерства науки и высшего образования России (специальности 12.00.01, 12.00.02), выходит 1 раз в месяц.
- Научно-правовой журнал «Право и жизнь», рецензируемый (РИНЦ, E-Library), выходит 1 раз в 3 месяца.