

КОНЦЕПЦИЯ «ПРЕДСКАЗУЕМОГО ВРЕДА» ПРИ РАЗРАБОТКЕ МЕДИЦИНСКИХ ИЗДЕЛИЙ НА ОСНОВЕ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА

И. Р. Бегишев [✉], А. А. Шутова

Казанский инновационный университет имени В. Г. Тимирязова, Казань, Россия

В статье представлена концепция «предсказуемого вреда» как методологический инструмент для комплексной оценки рисков при разработке и внедрении медицинских изделий на основе искусственного интеллекта. Актуальность исследования обусловлена экспоненциальным ростом применения ИИ-технологий в здравоохранении при одновременном отсутствии унифицированных подходов к прогнозированию потенциальных негативных последствий их использования. Проведен критический анализ существующих регуляторных подходов к оценке рисков, включая отечественные нормативные документы и международные стандарты. Предложена многомерная классификация типов предсказуемого вреда с учетом всего жизненного цикла медицинских ИИ-систем. Особое внимание уделено этическим аспектам применения искусственного интеллекта в медицине, включая принципы автономии пациента, справедливости, непричинения вреда и прозрачности алгоритмов. Разработана расширенная матрица оценки предсказуемого вреда, интегрирующая технологические, клинические и этические параметры для каждого этапа разработки и внедрения ИИ-систем в медицинскую практику. Результаты исследования могут служить методологической основой для разработчиков медицинских ИИ-систем, регуляторных органов и медицинских организаций при оценке безопасности и эффективности внедрения интеллектуальных технологий в клиническую практику.

Ключевые слова: искусственный интеллект, медицинские изделия, предсказуемый вред, этика искусственного интеллекта, регулирование, безопасность пациентов, управление рисками

Финансирование: работа выполнена за счет гранта Академии наук Республики Татарстан, предоставленного молодым кандидатам наук (постдокторантам) с целью защиты докторской диссертации, выполнения научно-исследовательских работ, а также выполнения трудовых функций в научных и образовательных организациях Республики Татарстан в рамках Государственной программы Республики Татарстан «Научно-технологическое развитие Республики Татарстан».

Благодарности: авторы выражают искреннюю признательность Академии наук Республики Татарстан за оказанную финансовую поддержку данного исследования по результатам конкурсного отбора (грант № 153/2024-ПД от 16 декабря 2024 г.) для научно-исследовательского проекта «Обеспечение технологического суверенитета системы здравоохранения уголовно-правовыми средствами», а также глубокую благодарность коллективу Научно-исследовательского института цифровых технологий и права Казанского инновационного университета имени В. Г. Тимирязова за ценные консультации и конструктивное обсуждение концептуальных положений исследования.

Вклад авторов: И. Р. Бегишев — концептуализация исследования; разработка теоретических основ понятия «предсказуемого вреда»; анализ специфики рисков, связанных с применением искусственного интеллекта в медицинских изделиях; систематизация типологии предсказуемого вреда; исследование нормативных правовых основ регулирования ИИ-систем в здравоохранении в различных юрисдикциях; формулирование выводов исследования; подготовка первоначального варианта рукописи; А. А. Шутова — разработка методологии прогнозирования и превенции предсказуемого вреда; создание матрицы оценки предсказуемого вреда на различных этапах жизненного цикла медицинских ИИ-систем; формирование рекомендаций по практической имплементации концепции предсказуемого вреда; анализ источников литературы; редактирование и критический пересмотр рукописи с внесением ценного интеллектуального содержания; визуализация исследования (разработка таблиц).

Соблюдение этических стандартов: заседание этического комитета не требовалось, поскольку данное исследование носит теоретико-методологический характер и основано на анализе открытых литературных источников и нормативных правовых документов, без проведения экспериментов с участием людей или животных и без использования персональных данных пациентов.

✉ **Для корреспонденции:** Ильдар Рустамович Бегишев
ул. Московская, д. 42, Республика Татарстан, г. Казань, 420111; begishev@mail.ru

Статья поступила: 03.05.2025 **Статья принята к печати:** 16.05.2025 **Опубликована онлайн:** 20.06.2025

DOI: 10.24075/medet.2025.006

THE CONCEPT OF PREDICTABLE HARM IN DEVELOPMENT OF AI-POWERED MEDICAL DEVICES

Begishev IR [✉], Shutova AA

Kazan Innovative University named after VG Timiryasov, Kazan, Russia

The article reviews the concept of predictable harm as a methodological tool for a comprehensive risk assessment while developing and implementing AI-powered medical devices. The study is relevant due to exponential growth of using AI-powered technologies in healthcare and lack of unified approaches to prediction of potential negative consequences of their usage. Existing regulatory approaches to risk assessment, including Russian regulatory documents and international standards, have been analyzed. A multidimensional classification of types of predictable harm is proposed considering the entire life cycle of medical AI systems. Special attention is given to ethical aspects of using artificial intelligence in medicine, including the principles of patient autonomy, equity, non-harm and transparency of algorithms. An expanded matrix for assessing predictable harm has been developed. It integrated technological, clinical and ethical parameters for each stage of development and implementation of AI systems in medical practice. The results of the study can be used as a methodological framework for developers of medical AI systems, regulatory authorities and medical organizations in assessing safety and effectiveness of introducing intelligent technologies into clinical practice.

Keywords: artificial intelligence, medical devices, predictable harm, ethics of artificial intelligence, regulation, patient safety, risk management

Financing: the work was carried out under the grant from the Academy of Sciences of the Republic of Tatarstan, provided to young candidates of sciences (postdoctoral fellows) so that they could defend their doctoral dissertation, perform research and labor functions in scientific and educational organizations of the Republic of Tatarstan within the framework of the Scientific and Technological Development of the Republic of Tatarstan State Program.

Acknowledgements: the authors express their sincere gratitude to the Academy of Sciences of the Republic of Tatarstan for financial support of this study based on the results of a competitive selection (grant No. 153/2024-PD dated December 16, 2024) for 'Ensuring the technological sovereignty of the healthcare system by criminal means' research project, as well as deep gratitude to the staff of the Scientific Research Institute of Digital Technologies and Law of Kazan Innovative University named after Timiryasov VG for valuable consultations and constructive discussion of the conceptual provisions of the study.

Author contribution: Begishev IR — research conceptualization; development of predictable harm theoretical foundations; analysis of specific risks associated with using artificial intelligence in medical devices; systematization of predictable harm typology; research of legal framework for regulating AI systems in healthcare with various jurisdictions; formulation of research conclusions; preparation of the manuscript initial version; Shutova AA — development of a methodology for predicting and preventing predictable harm; creation of a matrix for assessing predictable harm at various stages of medical AI system life cycle; developing recommendations for practical implementation of predictable harm concept; analysis of literature sources; editing and critical revision of the manuscript with introduction of valuable intellectual content; visualization of research (making tables).

Compliance with ethical standards: a meeting of the ethics committee was not required, since this study is theoretical and methodological in nature and analyzes open literature sources and regulatory legal documents, without conducting experiments involving humans or animals and without using personal patient data.

✉ **Correspondence should be addressed:** Ildar R. Begishev
42 Moskovskaya St., Kazan, Russia, 420111; begishev@mail.ru

Received: 03.05.2025 **Accepted:** 16.05.2025 **Published online:** 20.06.2025

DOI: 10.24075/medet.2025.006

Интеграция технологий искусственного интеллекта в сферу медицинских изделий открывает беспрецедентные возможности для совершенствования диагностических процессов, персонализации терапевтических подходов и оптимизации клинических решений. Однако стремительное внедрение ИИ-систем в здравоохранение сопряжено с возникновением специфических рисков, требующих систематического анализа и проактивного управления. В данном контексте концепция «предсказуемого вреда» приобретает ключевое значение как методологический инструмент превентивной идентификации и минимизации потенциальных негативных последствий применения медицинских изделий на основе искусственного интеллекта. Особую значимость при разработке данной концепции имеет критический анализ существующих регуляторных подходов к оценке рисков ИИ-систем в здравоохранении, а также интеграция этических принципов в процесс прогнозирования и предотвращения возможного вреда.

Актуальность исследования обусловлена экспоненциальным ростом рынка ИИ-решений в здравоохранении при одновременной недостаточности унифицированных подходов к оценке их безопасности. Согласно данным аналитического отчета Grand View Research, глобальный рынок искусственного интеллекта в медицине к 2028 г. достигнет 120,2 млрд долларов США с ежегодным приростом около 41,8%, что подчеркивает масштаб вызовов в области обеспечения безопасности пациентов [1].

Целью настоящего исследования является формирование методологических основ идентификации и минимизации предсказуемого вреда при разработке и имплементации медицинских изделий на основе искусственного интеллекта.

Для достижения поставленной цели определены следующие задачи.

1. Концептуализировать понятие предсказуемого вреда в контексте медицинских ИИ-технологий.
2. Проанализировать специфику рисков, ассоциированных с применением искусственного интеллекта в медицинских изделиях.
3. Исследовать существующие подходы к регулированию безопасности ИИ-систем в здравоохранении и провести сравнение с авторской концепцией.
4. Разработать методологию прогнозирования и превенции потенциального вреда при создании медицинских ИИ-систем с подробной этической проработкой.

ОСНОВНАЯ ЧАСТЬ

Концепция «предсказуемого вреда» в контексте медицинских изделий на основе искусственного интеллекта представляет собой методологический конструкт, интегрирующий принципы предиктивного анализа рисков, проактивного управления безопасностью и итеративной переоценки потенциальных негативных последствий применения технологии. Фундаментальным отличием данной концепции от традиционных подходов к оценке рисков является смещение фокуса с реактивного реагирования на инциденты к превентивному прогнозированию вероятных сценариев возникновения нежелательных явлений, обусловленных спецификой функционирования ИИ-систем.

Терминологическое определение предсказуемого вреда в рассматриваемом контексте может быть сформулировано как совокупность потенциальных негативных последствий применения медицинских ИИ-систем, которые могут быть идентифицированы и минимизированы на основе систематического анализа характеристик технологии, контекста ее применения и возможных траекторий эволюции системы в процессе эксплуатации. Ключевыми атрибутами данного определения являются: прогностический характер оценки, системный подход к анализу рисков и учет динамической природы ИИ-технологий.

На современном этапе уже существует ряд регуляторных подходов к оценке рисков при применении медицинских изделий на основе искусственного интеллекта.

Так, Приказ Министерства здравоохранения Российской Федерации от 7 июля 2020 г. № 686н [2] и письмо Федеральной службы по надзору в сфере здравоохранения от 13 февраля 2020 г. № 02И-297/20 [3] предусматривают ранжирование степени риска, согласно которому все медицинские изделия на основе искусственного интеллекта отнесены к 3-й степени риска еще до их применения и на этапе государственной регистрации. Этот подход ориентирован на централизованное регулирование и априорную высокорисковую классификацию всех ИИ-систем в медицине.

Международный форум регуляторов медицинских изделий (IMDRF, 2014) предлагает дифференцированную классификацию потенциальных рисков медицинских изделий на основе искусственного интеллекта в зависимости от клинического применения и возможного влияния на процесс лечения (Software as a Medical Device:

Possible Framework for Risk Categorization and Corresponding Considerations) [4]. Эта классификация учитывает и тяжесть потенциального вреда для пациента, и роль ИИ-системы в клиническом процессе.

Отраслевое приложение к Кодексу этики искусственного интеллекта Альянса в сфере искусственного интеллекта детализирует градацию рисков в зависимости от тяжести ошибок, связанных с применением систем искусственного интеллекта, акцентируя внимание на последствиях неверных врачебных решений [5].

Специфика рисков, ассоциированных с применением искусственного интеллекта в медицинских изделиях, обусловлена рядом уникальных характеристик данных технологий: автономностью функционирования, потенциальной непрозрачностью процесса принятия решений (проблема «черного ящика»), способностью к самообучению и адаптации, а также высокой зависимостью от качества исходных данных [6]. Данные особенности формируют многомерный профиль рисков, требующий дифференцированного подхода к их идентификации и управлению.

Предлагаемая концепция «предсказуемого вреда» отличается мультидименсиональной структурой, ориентированной на весь жизненный цикл ИИ-систем и включает:

- проактивную ориентацию на выявление рисков;
- дифференцированный подход к типам вреда (алгоритмический, датацентрический, интеграционный и др.);
- многоуровневую стратификацию ответственности участников;
- итеративность оценки рисков и адаптацию к эволюции ИИ-систем;
- интеграцию технологических и клинических аспектов качества.

В таблице 1 представлена систематизация типов предсказуемого вреда в контексте медицинских ИИ-систем.

Нормативное правовое регулирование безопасности медицинских изделий на основе искусственного интеллекта характеризуется значительной гетерогенностью подходов в различных юрисдикциях. Европейский союз имплементирует структурированную систему регулирования через Регламент о медицинских изделиях (MDR 2017/745) [7]

и Регламент об искусственном интеллекте (AI Act) [8], классифицирующий медицинские ИИ-системы как высокорисковые и устанавливающий строгие требования к их прозрачности и валидации.

Управление по санитарному надзору за качеством пищевых продуктов и медикаментов США (FDA) реализует адаптивный подход, основанный на концепции «Pre-Certification Program», фокусирующейся на оценке процессов разработки и культуры качества разработчика [9]. Данный подход предполагает непрерывный мониторинг производительности системы в реальных условиях эксплуатации и итеративную переоценку профиля риск—польза.

На основе анализа существующих подходов и регуляторных требований предлагается интегрированная методология прогнозирования и превенции предсказуемого вреда при разработке медицинских изделий на основе искусственного интеллекта, включающая следующие компоненты: многоуровневая модель оценки рисков, система инклюзивной валидации на разнородных популяциях пациентов, механизмы обеспечения интерпретируемости алгоритмов, инфраструктура непрерывного мониторинга производительности и процессы итеративной переоценки безопасности.

Практическая имплементация предложенной методологии может быть реализована через матрицу оценки предсказуемого вреда, представленную в таблице 2.

Этические аспекты являются неотъемлемой частью концепции «предсказуемого вреда» и отражаются на всех этапах жизненного цикла медицинских изделий на основе искусственного интеллекта. Рассмотрим основные измерения этической ответственности.

1. Автономия пациента — критически важно обеспечить, чтобы внедрение ИИ-систем не снижало роли пациента в процессе принятия решений. Следует учитывать риски чрезмерного автоматического доверия клиницистов к советам искусственного интеллекта и качество информированного согласия.
2. Справедливость — предотвращение алгоритмической дискриминации и обеспечение доступа различных групп населения к ИИ-технологиям. Необходимо акцентировать внимание на датацентрических рисках и репрезентативности данных.

Таблица 1. Типология предсказуемого вреда в контексте медицинских ИИ-систем

Категория вреда	Характеристика	Примеры	Методы идентификации и минимизации
Алгоритмический	Связан с дефектами в математических моделях и логике принятия решений	Ложноположительные / ложноотрицательные результаты, ошибки классификации	Валидация на репрезентативных выборках, тестирование граничных случаев
Датацентрический	Обусловлен проблемами в обучающих данных	Систематические смещения, неприменимость к определенным группам пациентов	Аудит данных, стратификация выборок, контроль репрезентативности
Интеграционный	Возникает при взаимодействии ИИ-системы с клиническими процессами	Нарушение протоколов, конфликты с существующими системами	Процессное моделирование, симуляция клинических сценариев
Интерпретационный	Связан с некорректной интерпретацией результатов пользователями	Переоценка/недооценка рекомендаций ИИ, пропуск критической информации	Оптимизация интерфейса, обучение пользователей
Эволюционный	Обусловлен изменением поведения системы в процессе самообучения	Дрейф концепций, деградация точности	Мониторинг производительности, периодическая реаттестация

Таблица 2. Матрица оценки предсказуемого вреда для медицинских ИИ-систем

Этап	Ключевые вопросы оценки	Инструменты	Ответственные стороны	Клинические аспекты	Этические аспекты
Концептуализация и дизайн	Соответствие целевому применению, охват клинических сценариев, техническая выполнимость	Этический аудит, анализ клинических сценариев, обзор доказательной базы	Разработчики, клинические эксперты, этические комитеты	Оценка клинической значимости решаемой задачи, потенциальных изменений в клинической практике, соотношения риск/польза	Соответствие ценностям медицинской профессии, обеспечение автономии пациента, соответствие принципам медицинской этики
Разработка и обучение	Репрезентативность данных, валидность алгоритма, устойчивость к экстремальным случаям	Статистический анализ, алгоритмический аудит, симуляция граничных случаев	Разработчики, специалисты по данным, ML-инженеры	Охват разнообразных клинических ситуаций, включение редких случаев, учет коморбидных состояний	Предотвращение дискриминации и смещений по признакам пола, возраста, этнической принадлежности, социально-экономического статуса
Валидация и верификация	Точность, специфичность, чувствительность, робастность, производительность	Кросс-валидация, тестирование граничных случаев, внешняя валидация	Независимые эксперты, клиницисты, регуляторы	Оценка влияния на клинические решения и исходы лечения, сравнение с «золотым стандартом» диагностики	Прозрачность и объяснимость результатов, возможность оспаривания, защита от автоматизированной дискриминации
Внедрение и интеграция	Совместимость с протоколами, влияние на клинические решения, удобство использования	Симуляция рабочих процессов, тестирование в реальных условиях, аудит клинических путей	Медицинские организации, IT-специалисты, клиницисты	Оценка влияния на процесс оказания помощи, время принятия решений, интеграция в существующие клинические протоколы	Влияние на взаимоотношения врач-пациент, уровень доверия, сохранение клинической автономии врача
Постмаркетинговый мониторинг	Нежелательные явления, дрейф производительности, непредвиденные последствия	Аналитика реальных данных, система сообщений о происшествиях, регулярный аудит	Производители, регуляторные органы, медицинские специалисты, пациентские сообщества	Мониторинг отклонений клинических исходов от ожидаемых, долгосрочное влияние на качество помощи, выявление редких осложнений	Учет опыта пациентов, психосоциальные последствия применения ИИ, оценка воздействия на доступность медицинской помощи

3. Непричинение вреда — учет возможности отложенных и системных последствий, связанных с эволюционными изменениями и самообучающимися алгоритмами.
4. Прозрачность и объяснимость — обеспечение интерпретируемости решений и возможности аудита как для специалистов, так и пациентов; преодоление эффекта «черного ящика».
5. Обязательный этический аудит — анализ соответствия использования искусственного интеллекта медицинской этике, регулярная ревизия матрицы риска с учетом уязвимости отдельных категорий пациентов и долгосрочных последствий.

Эти положения отражаются в матрице оценки предсказуемого вреда, представлены в расширенном виде (см. табл. 2).

Использование данной матрицы позволяет структурировать процесс идентификации и минимизации предсказуемого вреда на всех этапах жизненного цикла медицинского изделия на основе искусственного интеллекта, обеспечивая комплексный подход к управлению рисками и соответствие регуляторным и этическим требованиям.

Решающее значение для эффективной имплементации концепции предсказуемого вреда имеет формирование культуры прозрачности при разработке медицинских ИИ-систем. Данный аспект включает открытую коммуникацию относительно ограничений технологии, активное вовлечение клинических

специалистов на всех этапах создания продукта и применение принципа «безопасность через дизайн», предполагающего интеграцию механизмов обеспечения безопасности непосредственно в архитектуру системы. В отличие от существующих регуляторных подходов, фокусирующихся преимущественно на технических характеристиках и предварительной классификации рисков, предлагаемая концепция предполагает обязательную интеграцию этических аудитов на каждом этапе жизненного цикла медицинской ИИ-системы. Это требует мультидисциплинарного взаимодействия между разработчиками, клиницистами, специалистами по этике и представителями пациентских сообществ для предотвращения алгоритмической дискриминации, сохранения автономии пациента и поддержания справедливого доступа к преимуществам технологий. Матрица оценки предсказуемого вреда с включением этических и клинических параметров становится не просто инструментом документирования, но и платформой для непрерывного диалога между всеми участниками процесса внедрения ИИ в клиническую практику.

ВЫВОДЫ И ЗАКЛЮЧЕНИЕ

Проведенное исследование позволяет сформулировать следующие основные выводы.

1. Концепция «предсказуемого вреда» — эффективный инструмент для повышения безопасности внедрения

искусственного интеллекта в медицину, превентивно выявляющий и минимизирующий риски.

2. Уникальные риски применения искусственного интеллекта требуют дифференцированного подхода к управлению — с обязательной интеграцией этических аспектов.
3. Сравнительный анализ показал, что авторская концепция дополняет и расширяет существующие регуляторные подходы, обеспечивая многомерную, непрерывную и этически подкрепленную оценку вреда.
4. Перспективы дальнейшей работы связаны с универсализацией методологий и стандартизацией практик оценки рисков создания и применения искусственного интеллекта в здравоохранении.

Таким образом, развитие и практическая имплементация концепции предсказуемого вреда в процессы разработки и внедрения медицинских изделий на основе искусственного интеллекта представляются необходимым условием для обеспечения оптимального баланса между инновационным потенциалом данных

технологий и безопасностью пациентов. Сравнительный анализ с существующими регуляторными подходами демонстрирует преимущества предлагаемой концепции в аспекте мультидименсиональной оценки рисков и интеграции этических принципов на всех стадиях жизненного цикла ИИ-систем. Матрица предсказуемого вреда, включающая параметры клинических последствий и этической оценки, позволяет перейти от формальных процедур управления рисками к системному подходу, учитывающему как технологические, так и гуманитарные аспекты применения искусственного интеллекта в медицине. Перспективными направлениями дальнейших исследований в данной области являются разработка стандартизированных методологий оценки рисков для различных категорий ИИ-систем, создание валидированных инструментов этического аудита медицинских ИИ-решений и формирование унифицированных регуляторных требований, синтезирующих технологические стандарты с принципами медицинской этики и ориентированных на долгосрочные социальные последствия внедрения интеллектуальных технологий в здравоохранении.

Литература

1. Grand View Research. Artificial Intelligence in Healthcare Market Size Report, 2021–2028. Available from URL: <https://www.grandviewresearch.com/industry-analysis/artificial-intelligence-ai-healthcare-market> (accessed: 01.05.2025).
2. Приказ Министерства здравоохранения Российской Федерации от 7 июля 2020 г. № 686н «О внесении изменений в приложения № 1 и № 2 к приказу Министерства здравоохранения Российской Федерации от 6 июня 2012 г. № 4н «Об утверждении номенклатурной классификации медицинских изделий» Официальный интернет-портал правовой информации. Электрон. дан. Режим доступа: [Электронный ресурс] URL: <https://www.pravo.gov.ru>, свободный. № опубликования: 0001202008100015 (дата обращения: 10.05.2025).
3. Письмо Федеральной службы по надзору в сфере здравоохранения от 13 февраля 2020 г. № 02И-297/20 «О программном обеспечении». [Текст письма опубликован не был].
4. «Software as a Medical Device»: Possible Framework for Risk Categorization and Corresponding Considerations. Available from URL: <https://www.imdrf.org/sites/default/files/docs/imdrf/final/technical/imdrf-tech-140918-samd-framework-risk-categorization-141013.pdf> (accessed: 10.05.2025).
5. Кодекс этики в сфере ИИ в медицине и здравоохранении. Режим доступа: [Электронный ресурс] URL: <https://ethics.a-ai.ru/ethics-of-medicine> (accessed: 10.05.2025).
6. World Health Organization. Ethics and governance of artificial intelligence for health: WHO guidance. Available from URL: <https://www.who.int/publications/i/item/9789240029200> (accessed: 01.05.2025).
7. European Parliament. Regulation (EU) 2017/745 of the European Parliament and of the Council of 5 April 2017 on medical devices, amending Directive 2001/83/EC, Regulation (EC) No 178/2002 and Regulation (EC) No 1223/2009 and repealing Council Directives 90/385/EEC and 93/42/EEC (Text with EEA relevance.). Available from URL: <https://eur-lex.europa.eu/eli/reg/2017/745/oj/eng> (accessed: 01.05.2025).
8. European Commission. Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act). Available from URL: <https://digital-strategy.ec.europa.eu/en/library/proposal-regulation-laying-down-harmonised-rules-artificial-intelligence> (accessed: 01.05.2025).
9. U. S. Food and Drug Administration. Artificial Intelligence and Machine Learning (AI/ML) Software as a Medical Device Action Plan. Available from URL: <https://www.fda.gov/medical-devices/software-medical-device-samd/artificial-intelligence-and-machine-learning-software-medical-device> (accessed: 01.05.2025).

References

1. Grand View Research. Artificial Intelligence in Healthcare Market Size Report, 2021–2028. Available from URL: <https://www.grandviewresearch.com/industry-analysis/artificial-intelligence-ai-healthcare-market> (accessed: 01.05.2025).
2. Prikaz Ministerstva zdravookhraneniya Rossiyskoy Federatsii ot 7 iyulya 2020 g. № 686n «O vnesenii izmeneniy v prilozheniya № 1 i № 2 k prikazu Ministerstva zdravookhraneniya Rossiyskoy Federatsii ot 6 iyunya 2012 g. № 4n «Ob utverzhenii nomenklturnoy klassifikatsii meditsinskikh izdeliy» Ofitsial'nyy internet-portal zakona informatsii. Elektron. dan. Available from URL: <https://www.pravo.gov.ru>, svobodnyy. № publikatsii: 0001202008100015 (accessed: 10.05.2025). Russian.
3. Pis'mo Federal'noy sluzhby po nadzoru v sfere zdravookhraneniya ot 13 fevralya 2020 g. № 02I-297/20 «O programmnom obespechenii». [Tekst pis'ma opublikovan ne byl]. Russian.
4. «Software as a Medical Device»: Possible Framework for Risk Categorization and Corresponding Considerations. Available from URL: <https://www.imdrf.org/sites/default/files/docs/imdrf/final/technical/imdrf-tech-140918-samd-framework-risk-categorization-141013.pdf> (accessed: 10.05.2025).
5. Kodeks etiki v sfere II v meditsine i zdravookhraneni. Available from URL: <https://ethics.a-ai.ru/ethics-of-medicine> (accessed: 10.05.2025). Russian.
6. World Health Organization. Ethics and governance of artificial intelligence for health: WHO guidance. Available from URL: <https://www.who.int/publications/i/item/9789240029200>

- www.who.int/publications/i/item/9789240029200 (accessed: 01.05.2025).
7. European Parliament. Regulation (EU) 2017/745 of the European Parliament and of the Council of 5 April 2017 on medical devices, amending Directive 2001/83/EC, Regulation (EC) No 178/2002 and Regulation (EC) No 1223/2009 and repealing Council Directives 90/385/EEC and 93/42/EEC (Text with EEA relevance.). Available from URL: <https://eur-lex.europa.eu/eli/reg/2017/745/oj/eng> (accessed: 01.05.2025).
 8. European Commission. Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act). Available from URL: <https://digital-strategy.ec.europa.eu/en/library/proposal-regulation-laying-down-harmonised-rules-artificial-intelligence> (accessed: 01.05.2025).
 9. U. S. Food and Drug Administration. Artificial Intelligence and Machine Learning (AI/ML) Software as a Medical Device Action Plan. Available from URL: <https://www.fda.gov/medical-devices/software-medical-device-samd/artificial-intelligence-and-machine-learning-software-medical-device> (accessed: 01.05.2025).