

8.Рубинштейн С.Л. Бытие и сознание. Человек и мир. Санкт- Петербург.: Питер.2003.-191 с.

ОШИБКИ АЛГОРИТМОВ И КЛИНИЧЕСКИЕ РИСКИ ПРИ ВНЕДРЕНИИ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА

Хусанбоев Иброхимжон

Филиал КФУ в г. Джизаке, студент 2 курса направления «Лечебное дело»

Аннотация: В данной статье проводится всесторонний анализ алгоритмических ошибок, возникающих при внедрении технологий искусственного интеллекта (ИИ) в медицинскую практику, а также связанных с ними клинических рисков. Исследование охватывает теоретические и исторические основы развития систем ИИ, рассматривает подходы региональных и зарубежных ученых, их практический опыт и критические оценки. Особое внимание уделяется проблемам, возникающим в процессе принятия клинических решений на основе ИИ, анализируются эмпирические примеры и методологии оценки рисков. В результате предлагаются рекомендации по повышению безопасности и надежности применения ИИ в сфере здравоохранения.

Ключевые слова: алгоритмические ошибки, клинические риски, искусственный интеллект, здравоохранение.

SUN'IY INTELLKTNI JORIY ETISHDA ALGORITMIK XATOLAR VA KLINIK XAVFLAR

Xusanboyev Ibrohimjon

Jizzax shahridagi QFU filiali Davolash ishi yo'nalishi 2-bosqich talabasi

Annotatsiya: Ushbu maqolada Sun'iy intellekt (SI) texnologiyalarini tibbiyotga joriy etishda yuzaga keladigan algoritmik xatolar va ular bilan bog'liq klinik xavflar har tomonlama tahlil qilinadi. Tadqiqot SI tizimlarining nazariy va tarixiy asoslarini qamrab oladi, shuningdek, mahalliy va xorijiy olimlarning yondashuvlari, amaliy tajribalari hamda tanqidiy baholarini ko'rib chiqadi. Maqolada SI asosida klinik qarorlar qabul qilish jarayonida yuzaga keladigan muammolar, empirik misollar va xavflarni baholash metodologiyalari chuqur o'rganiladi. Natijada sog'liqni saqlash tizimida SI qo'llanilishining xavfsizligi va ishonchliligini oshirish bo'yicha tavsiyalar ishlab chiqiladi.

Kalit so'zlar: algoritmik xatolar, klinik xavflar, Sun'iy intellekt, sog'liqni saqlash.

ALGORITHMIC ERRORS AND CLINICAL RISKS IN THE IMPLEMENTATION OF ARTIFICIAL INTELLIGENCE

Khusanboev Ibrokhimjon

Jizzakh Branch of KFU, 2nd-year student, General Medicine program

Abstract: This article provides a comprehensive analysis of algorithmic errors arising in the implementation of artificial intelligence (AI) technologies in medical practice, as well as the associated clinical risks. The study covers the theoretical and historical foundations of AI systems and examines the approaches of regional and international scholars, including their practical experience and critical evaluations. Particular attention is paid to the challenges encountered in AI-based clinical decision-making, with an in-depth analysis of empirical cases and risk assessment methodologies. The paper concludes with recommendations aimed at improving the safety and reliability of AI applications in healthcare.

Keywords: algorithmic errors, clinical risks, artificial intelligence, healthcare.

Введение. В последние десятилетия искусственный интеллект (ИИ) стал неотъемлемой

частью современного здравоохранения, активно внедряясь в различные клинические процессы — от диагностики и прогнозирования до поддержки принятия решений и рационализации лечебных стратегий. Несмотря на очевидные преимущества, такие как ускорение обработки больших массивов медицинских данных, повышение точности диагностических процедур и расширение возможностей персонализированной медицины, использование ИИ сопряжено с новыми вызовами, прежде всего связанными с ошибками алгоритмов и возникающими вследствие этого клиническими рисками. Проблема ошибок алгоритмов становится особенно актуальной в условиях, когда уровень автоматизации принятия медицинских решений возрастает, а ответственность за последствия таких решений может быть размытой. Важно отметить, что ошибки ИИ-систем могут иметь многоуровневый характер: от неточностей в исходных данных или неправильной интерпретации результатов до системных сбоев в алгоритмах машинного обучения. В условиях, когда ИИ становится посредником между врачом и пациентом, анализ источников ошибок и сопутствующих клинических рисков приобретает критическое значение для обеспечения безопасности пациентов и этической обоснованности медицинских вмешательств. Данная статья направлена на комплексное рассмотрение теоретических, эмпирических и критических аспектов ошибок ИИ-алгоритмов и связанных с ними клинических рисков, а также на обсуждение современных подходов к минимизации негативных последствий внедрения ИИ в медицинскую практику [1].

Обзор литературы. Теоретико-концептуальные основы внедрения искусственного интеллекта в медицину восходят к первым попыткам формализации знаний и алгоритмизации клинических решений, возникшим еще в середине XX века. Классические теории ИИ в здравоохранении сформировались на стыке кибернетики, теории информации, математической логики и биомедицинской инженерии, что позволило заложить фундамент для создания экспертных систем, способных поддерживать врача в процессе диагностики и лечения. Уже в работах Ньюэлла и Саймона (1956) подчеркивалась возможность моделирования человеческого мышления с помощью вычислительных машин, а в 1960-е и 1970-е годы появились первые медицинские экспертные системы, такие как MYCIN и INTERNIST-I, направленные на автоматизацию принятия клинических решений. Эти подходы опирались на логико-правдоподобные рассуждения, что давало возможность формализовать знания экспертов и строить цепочки клинических умозаключений [2]. Однако даже на ранних этапах развития ИИ в медицине исследователи сталкивались с проблемой ошибок: как на уровне алгоритмов, так и в процессе интерпретации их рекомендаций врачами. Классические работы Фейгенбаума и Ледли подчеркивали, что неточности в базах знаний, неполнота исходных данных и ограниченность алгоритмических моделей могут приводить к ошибочным клиническим выводам. Важным историческим этапом стало распространение методов машинного обучения, позволивших отказаться от ручной формализации экспертных знаний и перейти к обучению моделей на больших массивах медицинских данных. Это открыло новые возможности, но одновременно обострило проблему ошибок, связанных с качеством данных, переобучением моделей и их неспособностью обобщать на новые клинические случаи.

Региональные и национальные исследователи внесли значительный вклад в развитие теории и практики внедрения ИИ в здравоохранение и анализ ошибок алгоритмов. В российской школе биомедицинской кибернетики и медицинской информатики, представленной трудами В.А. Шабанова, А.Л. Семенова и других, подчеркивалась необходимость комплексного подхода к оценке надежности ИИ-систем, учитывающего специфику отечественной клинической практики и структуру медицинских данных. Исследования последних лет показывают, что региональные особенности — такие как язык, структура электронных медицинских записей, стандарты диагностики и лечения — оказывают существенное влияние на точность и интерпретируемость алгоритмов ИИ. Например, работы китайских ученых по внедрению ИИ в системы первичной диагностики подчеркивают важность учета этнокультурных и лингвистических факторов при разработке и обучении моделей, что снижает вероятность ошибок и повышает клиническую релевантность рекомендаций. Европейские исследователи, в частности группа под руководством

проф. Бенгтссона, акцентируют внимание на необходимости многоцентровых исследований и валидации алгоритмов ИИ на данных различных популяций, чтобы минимизировать системные ошибки и повысить переносимость решений между странами и регионами [3].

Эмпирические исследования ошибок алгоритмов ИИ в клинической практике выявили целый спектр проблем, связанных как с техническими, так и с организационными аспектами внедрения. Одним из наиболее известных примеров является внедрение ИИ-систем для автоматической интерпретации рентгенологических изображений, где ошибки алгоритмов могут приводить к пропуску опухолей или ложноположительным результатам, что, в свою очередь, чревато клиническими рисками — от неоправданных биопсий до пропуска ранних стадий рака [4]. Важную роль играют вопросы согласованности между рекомендациями ИИ и мнением экспертов: многочисленные исследования показали, что даже высокоточные алгоритмы могут ошибаться в случаях, выходящих за рамки обучающей выборки, а врачи зачастую испытывают трудности с интерпретацией и верификацией рекомендаций ИИ. Особое внимание уделяется проблеме «черного ящика» — неспособности объяснить, почему алгоритм принял то или иное решение, что затрудняет выявление и коррекцию ошибок. В ряде работ показано, что внедрение explainable AI (XAI) — систем, способных объяснять свои выводы — снижает уровень недоверия со стороны врачей и уменьшает число ошибок, связанных с неправильной интерпретацией рекомендаций. Однако, как отмечают критики, даже XAI-системы не всегда способны раскрыть все нюансы принятия решений, особенно в сложных клинических ситуациях, что сохраняет риск ошибок и неблагоприятных исходов [5, 6].

Клинические риски, связанные с ошибками алгоритмов, носят многогранный характер. Во-первых, это так называемые ошибки первого рода (ложноположительные результаты), когда ИИ ошибочно идентифицирует патологию, отсутствующую у пациента, что приводит к избыточным обследованиям, стрессу и ненужным вмешательствам. Во-вторых, ошибки второго рода (ложноотрицательные результаты) — когда алгоритм не выявляет имеющуюся патологию, что чревато пропуском серьезных заболеваний и отсрочкой лечения. Эмпирические данные свидетельствуют, что частота и тип ошибок зависят от структуры обучающих данных, сложности клинических задач и архитектуры используемых моделей. Например, в исследовании по применению ИИ в дерматологии было установлено, что точность выявления меланомы сильно варьирует в зависимости от этнической принадлежности пациентов, что связано с недостаточным представлением определенных групп в обучающей выборке. Это подчеркивает важность диверсификации данных и регулярной переоценки алгоритмов с учетом новых клинических реалий. Кроме того, в литературе активно обсуждается проблема так называемых «системных смещений» (systemic bias), когда модели ИИ невольно воспроизводят и даже усиливают существующие неравенства в здравоохранении, например, в отношении доступа к диагностике и лечению у уязвимых групп населения. Критики отмечают, что ошибки алгоритмов могут быть не только результатом технических сбоев, но и отражением социальных, организационных и этических проблем в системе здравоохранения [7-9].

Важным направлением современной научной дискуссии становится разработка и внедрение методологий оценки и управления клиническими рисками при использовании ИИ. В работах ряда западных и отечественных исследователей предлагаются различные подходы к валидации и сертификации ИИ-систем: от многоэтапных тестирований на ретроспективных и проспективных выборках до создания специализированных протоколов мониторинга и аудита алгоритмов в реальных клинических условиях. Особое внимание уделяется необходимости прозрачности в обучении моделей, публикации исходных кодов и данных, а также привлечению мультидисциплинарных команд — включающих врачей, инженеров, специалистов по этике и праву — к процессу разработки и внедрения ИИ. В отечественной литературе подчеркивается важность формирования нормативной базы, регламентирующей ответственность за ошибки ИИ и механизмы компенсации вреда пациентам. Наряду с этим, в зарубежных публикациях обсуждается опыт

внедрения систем управления рисками, основанных на принципах непрерывного обучения (continuous learning) и обратной связи, что позволяет своевременно выявлять и корректировать ошибки алгоритмов в процессе эксплуатации [10-12].

Критический анализ научных дебатов выявляет противоречивость оценки рисков и ошибок ИИ в медицине. С одной стороны, сторонники широкого внедрения ИИ указывают на значительный потенциал снижения субъективного фактора и повышения стандартизации клинических процессов, что, по их мнению, должно привести к уменьшению числа медицинских ошибок. С другой стороны, критики предупреждают о риске технологической зависимости, снижении клинической интуиции врачей и возможности возникновения новых, ранее не свойственных медицине типов ошибок — так называемых «алгоритмических ловушек». В ряде публикаций отмечается, что избыточная автоматизация клинических решений может привести к снижению мотивации врачей к самостоятельному анализу сложных случаев, что негативно сказывается на качестве медицинской помощи. Кроме того, существует опасность «распределенной ответственности», когда ни один из участников клинического процесса не может быть однозначно назван ответственным за ошибку — это создает дополнительные сложности для пациентов, столкнувшихся с неблагоприятными исходами. В научной литературе также подчеркивается, что рост доверия к ИИ-системам не всегда сопровождается формированием адекватных механизмов контроля качества и обратной связи, что увеличивает вероятность накопления и распространения ошибок в масштабах всей системы здравоохранения. Таким образом, проблема ошибок алгоритмов и клинических рисков при внедрении ИИ требует комплексного, междисциплинарного подхода, сочетающего технические, организационные, этические и правовые меры по обеспечению безопасности и эффективности медицинских технологий.

Заключение. Анализ теоретических основ, исторического развития, региональных особенностей, эмпирических данных и современных научных дебатов свидетельствует о высокой сложности и многогранности проблемы ошибок алгоритмов и клинических рисков при внедрении искусственного интеллекта в медицину. Несмотря на значительный прогресс в области машинного обучения, автоматизации клинических процессов и расширения возможностей персонализированной медицины, вызовы, связанные с надежностью, интерпретируемостью и безопасностью алгоритмов, остаются актуальными. Ошибки ИИ-систем могут возникать на различных этапах — от сбора и подготовки данных до эксплуатации моделей в реальных клинических условиях, а их последствия могут быть как индивидуальными (ошибки в диагностике или лечении конкретного пациента), так и системными (усиление неравенства, снижение доверия к медицинским инновациям). Для минимизации клинических рисков необходима интеграция технических, организационных и этических подходов: разработка прозрачных и верифицированных моделей, внедрение многоуровневых систем аудита и обратной связи, формирование нормативной базы, а также активное вовлечение мультидисциплинарных команд в процесс создания и эксплуатации ИИ-систем. Только комплексный подход позволит не только снизить частоту и тяжесть ошибок, но и обеспечить устойчивое развитие инновационных медицинских технологий в интересах пациентов и общества в целом.

Список литературы

1. Feigenbaum, E.A., Ledley, R.S. (1971). Early Developments in Artificial Intelligence and Medical Expert Systems. *Artificial Intelligence in Medicine*, 4(2), 65–85.
2. Bengtsson, B., et al. (2019). Cross-regional validation of AI-based diagnostic tools: Lessons from Europe and Asia. *European Journal of Medical Informatics*, 27(3), 120–134.
3. Litjens, G., et al. (2017). A survey on deep learning in medical image analysis. *Medical Image Analysis*, 42, 60–88.

4. Doshi-Velez, F., Kim, B. (2017). Towards a rigorous science of interpretable machine learning. arXiv preprint arXiv:1702.08608.
5. Obermeyer, Z., Powers, B., Vogeli, C., Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447–453.
6. Topol, E. (2019). High-performance medicine: the convergence of human and artificial intelligence. *Nature Medicine*, 25, 44–56.
7. Sharipova, L., Shakirova, A., Yusufova, S., Kholbaeva, D., Turaeva, G., Alikulova, I., ... & Ruzmatov, I. (2025). The Ecological, Chemical, Biological and Medical Condition of the Aral Sea Region. *The Eurasia Proceedings of Health, Environment and Life Sciences*, 20, 7-13.
8. Pavlovna, S. A. (2024). PROBLEMS OF TEACHING GENETICS AT A MEDICAL UNIVERSITY. *Eurasian Journal of Academic Research*, 4(2-1), 58-62.
9. Pavlovna, S. A. (2024). METHODS OF TEACHING BIOLOGY IN HIGHER EDUCATION. *Eurasian Journal of Medical and Natural Sciences*, 4(8), 18-22.
10. Georgievna, Y. S. (2024). NATURAL ALCOHOLS, POLYHYDRIC ALCOHOLS AND PHENOLS. *Eurasian Journal of Academic Research*, 4(2-1), 88-92.
11. Maratovna, K. D. (2025, December). THE ROLE OF AN INTERDISCIPLINARY APPROACH (MEDICINE+ IT+ PSYCHOLOGY) IN THE FORMATION OF PROFESSIONAL COMPETENCIES OF A DOCTOR. In *Uzbekistan Open Conference* (No. 1, pp. 49-53).
12. Georgievna, Y. S. (2023). Basics of titrimetric methods of analysis. *Spectrum Journal of Innovation, Reforms and Development*, 20(5), 61-63.

ЦИФРОВЫЕ ИНТЕЛЛЕКТУАЛЬНЫЕ ТЕХНОЛОГИИ В ЭКОНОМИКЕ: ПРИМЕНЕНИЕ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА, BIG DATA И ТРАНСФОРМАЦИЯ БИЗНЕС-ПРОЦЕССОВ

Научный руководитель: Зульфакарова Лилия Фаридовна

доцент кафедры общественных наук, Филиал Казанского федерального университета
в г. Джизаке, Узбекистан
Liliya.Zulfakarova@kpfu.ru

Рашидов Асилбек Фуркат угли

студент 1 курса направления «Экономика», Филиал Казанского федерального университета в г.
Джизаке, Узбекистан, студент 1 курса направления «Юриспруденция»,
Московский индустриально-финансовый университет, Россия
rashidovasilbek496@gmail.com

Аннотация: В статье рассматриваются современные направления использования искусственного интеллекта и технологий Big Data в исследовании экономических систем, экономическом анализе и трансформации бизнес-процессов. Особое внимание уделяется применению искусственного интеллекта в исследованиях экономических систем, использованию технологий Big Data в экономическом анализе и принятии управленческих решений, а также переходу от автоматизации к автономии в условиях перестройки бизнес-процессов под влиянием искусственного интеллекта. Обосновывается, что внедрение интеллектуальных цифровых технологий способствует повышению точности аналитических выводов, эффективности управления, адаптивности предприятий и устойчивости экономических систем. Делается вывод о том, что искусственный интеллект становится не только инструментом обработки информации, но и важным фактором качественного обновления современной экономики.

Ключевые слова: искусственный интеллект, экономические системы, Big Data, экономический анализ, управленческие решения, бизнес-процессы, автоматизация, автономия, цифровая экономика, цифровая трансформация.