

НАУЧНЫЙ ОБЗОР

DOI: <https://doi.org/10.17816/socm622965>

## Проблемы медицинского искусственного интеллекта. Часть 2

В.А. Бердутин<sup>1</sup>, Т.Е. Романова<sup>2</sup>, С.В. Романов<sup>3</sup>, О.П. Абаева<sup>1</sup>

<sup>1</sup> Государственный научный центр Российской Федерации — Федеральный медицинский биофизический центр имени А.И. Бурназяна, Москва, Россия;

<sup>2</sup> Приволжский исследовательский медицинский университет, Нижний Новгород, Россия;

<sup>3</sup> Приволжский окружной медицинский центр, Нижний Новгород, Россия

### АННОТАЦИЯ

Возможности искусственного интеллекта (ИИ) и машинного обучения растут беспрецедентными темпами. Эти технологии имеют множество полезных применений: от машинного перевода до анализа медицинских изображений.

В настоящее время разрабатывается множество таких приложений, а в долгосрочной перспективе ожидается лавинообразное нарастание их числа. К сожалению, слабостям и иным неприятным сторонам ИИ уделяется недостаточно внимания. В данном обзоре мы рассматриваем целый спектр уже известных проблем и возможных рисков, связанных с использованием инновационных нейросетевых технологий, обращая особое внимание на способы предотвращения реальных опасностей и потенциальных угроз с целью расширить круг заинтересованных лиц и профильных экспертов, участвующих в обсуждении актуальных вопросов кибербезопасности медицинского ИИ, формирования ответственного подхода к уязвимостям нейросетевых платформ, повышения надёжности защиты оборудования для его безопасного использования, а также к важности правовых и этических аспектов регулирования применения ИИ.

Несмотря на отдельные проблемы, описанные в нашем обзоре, очевидно, что ИИ будет важным элементом будущего здравоохранения. Поскольку население продолжает стареть, а спрос на медицинские услуги растёт, ожидается, что нейронные сети совсем скоро будут выступать в роли движущей силы здравоохранения, особенно в областях анализа медицинских изображений, виртуальных помощников, разработки лекарств, рекомендаций по лечению и обработки данных пациентов. Мы хотели бы подчеркнуть, что, признавая инновационную роль, которую цифровые технологии и ИИ могут и должны играть в укреплении отечественной системы здравоохранения, не стоит упускать из виду, насколько важно своевременно и правильно оценивать их благоприятное или негативное влияние на отрасль, чтобы обеспечить такие управленческие решения, которые бы неоправданно не отвлекали наше внимание и ресурсы от нецифровых подходов и исследований.

Настоящая статья представляет собой продолжение статьи: Бердутин В.А., Романова Т.Е., Романов С.В., Абаева О.П. Проблемы медицинского искусственного интеллекта. Часть 1 // Социология медицины. 2023. Т. 22, № 2. С. 202–211. DOI: <https://doi.org/10.17816/socm619132>

**Ключевые слова:** искусственный интеллект; машинное обучение; нейронная сеть.

### Как цитировать:

Бердутин В.А., Романова Т.Е., Романов С.В., Абаева О.П. Проблемы медицинского искусственного интеллекта. Часть 2 // Социология медицины. 2024. Т. 23, № 1. С. 94–103. DOI: <https://doi.org/10.17816/socm622965>

## REVIEW

DOI: <https://doi.org/10.17816/socm622965>

# Problematic aspects of medical artificial intelligence. Part 2

Vitalii A. Berdutin<sup>1</sup>, Tatyana E. Romanova<sup>2</sup>, Sergey V. Romanov<sup>3</sup>, Olga P. Abaeva<sup>1</sup><sup>1</sup> State Research Center — Burnasyan Federal Medical Biophysical Center, Moscow, Russia;<sup>2</sup> Privolzhsky Research Medical University, Nizhny Novgorod, Russia;<sup>3</sup> Privolzhsky District Medical Center, Nizhny Novgorod, Russia**ABSTRACT**

The capabilities of artificial intelligence (AI) and machine learning are growing at an unprecedented pace. These technologies have many useful applications, from machine translation to medical image analysis.

A large number of such applications are currently being developed, and an increasing number of such applications is expected in the long term. Unfortunately, weaknesses and other unpleasant aspects of AI have received insufficient attention. In this review, we consider a whole range of already known problems and possible risks associated with the use of innovative neural network technologies, paying special attention to the ways of preventing real dangers and potential threats in order to expand the range of stakeholders and profile experts involved in the discussion of current issues of medical AI cybersecurity, formation of responsible approach to the vulnerabilities of neural network platforms, increasing the reliability of equipment protection for its safe use, as well as the importance of legal and ethical aspects of regulating the use of AI.

Despite certain challenges described in our review, it is clear that AI will be an important element of the healthcare future. As the population continues to age and the demand for healthcare services grows, neural networks are expected to drive healthcare very soon, especially in the areas of medical image analysis, virtual assistants, drug development, treatment recommendations and patients' data processing. We would like to emphasize that while we recognize the innovative role that digital technologies and AI can and should play in strengthening the domestic healthcare system, we should not overlook the importance of timely and accurate assessment of their beneficial or negative impact on the industry to ensure such management decisions do not unnecessarily divert our attention and resources from non-digital approaches and research.

This article is a continuation of the article by Berdutin VA, Romanova TE, Romanov SV, Abaeva OP. Problematic aspects of medical artificial intelligence. Part 1. *Sociology of Medicine*. 2023;22(2):202–211. DOI: <https://doi.org/10.17816/socm619132>

**Keywords:** artificial intelligence; machine learning; neural network.

**To cite this article:**

Berdutin VA, Romanova TE, Romanov SV, Abaeva OP. Problematic aspects of medical artificial intelligence. Part 2. *Sociology of Medicine*. 2024;23(1):94–103. DOI: <https://doi.org/10.17816/socm622965>

## ОБОСНОВАНИЕ

Здравоохранение и медицинская практика претерпели значительные изменения в последние годы благодаря технологиям искусственного интеллекта (ИИ). Платформы ИИ открывают широкие перспективы, которые становятся известны медицинскому сообществу благодаря многочисленным научным публикациям, а также внедряемым повсеместно компьютерным приложениям и гаджетам. Данная публикация продолжает серию наших статей, посвящённых перспективному и проблемным аспектам использования систем ИИ в здравоохранении — от вопросов безопасного хранения персональной информации пациентов до рисков, связанных с развитием роботизированной хирургии. Мы стремимся привлечь внимание медицинского сообщества к ряду слабых сторон ИИ, возникающих в связи с использованием новых нейросетевых платформ. Кроме того, мы освещаем наиболее острые проблемы, описывая их потенциальное влияние и вызовы для системы здравоохранения. К одной из трудно решаемых проблем относится несанкционированный доступ к медицинским базам данных, который грозит негативными экономическими, психологическими и репутационными издержками. Сюда же можно отнести сложности со структурированием и нерелевантностью медицинской информации, поддержанием стабильности и робастности удалённого управления лечебно-диагностическим процессом, устранением аппаратных сбоев, а также с ускоренным решением задачи воспитания информационной грамотности медицинского персонала и т. д.

За последние 10-летия значительно возросло внимание к этическим аспектам применения медицинского ИИ. Среди прочего, активно обсуждаются экзистенциальные риски, связанные с дальнейшим развитием общего или *сильного* искусственного интеллекта — пока лишь гипотетической, но крайне опасной формы ИИ, способной на гораздо более быстрые и сложные действия, чем человек. Это привело к активным исследованиям того, как человечество может избежать потери контроля над ИИ, который превосходит по интеллекту лучших из нас. Отрадно, что сейчас начали активно разрабатываться *дружественные* платформы медицинского ИИ, т. е. гарантированно невраждебная для людей система, приоритетом которой является этико-деонтологическая направленность ИИ и личностные ценности пациента. В наших публикациях мы кратко рассматриваем ряд конкретных вопросов, связанных с использованием ИИ и машинного обучения (МО) в медицине, таких как справедливость, конфиденциальность, анонимность, интерпретируемость, а также более широкие социальные аспекты, включая этику и законодательство. Все эти аспекты необходимо учитывать для повышения общественного принятия технологий ИИ, обеспечения их соответствия развивающейся нормативно-правовой базе цифрового здравоохранения [1–6].

## О НЕУДАЧАХ МЕДИЦИНСКОГО ИСКУССТВЕННОГО ИНТЕЛЛЕКТА

Специалисты в области информационных технологий прекрасно осведомлены о проблеме снижения производительности платформы ИИ после её развертывания в реальных условиях. По очевидным причинам мелкие неудачи и даже крупные провалы разработчики стараются не афишировать. Однако публичные скандалы с гигантами ИТ-индустрии замолчать трудно, и они уже не единожды становились достоянием гласности. Один из самых громких произошёл, когда компания Google Verily Health Sciences провела полевые испытания своей системы для выявления диабетической ретинопатии в Таиланде. Как исследователи описали в академической статье [7], система работала плохо из-за недостаточного освещения и наличия большого количества изображений с низким разрешением: 21% изображений, которые техники пытались ввести, были отклонены моделью как неподходящие. Для остальных изображений авторы не раскрывают показатели точности, но говорят, что производительность заметно снизилась. Системе также часто требовалось много времени для запуска, поскольку изображения приходилось загружать в облако, что уменьшало количество людей, которых клиника могла обрабатывать каждый день.

Обнаружение рака кожи с помощью смартфона — одно из самых многообещающих направлений применения ИИ. Тем не менее каждая тестируемая сегодня система обнаружения рака кожи не может избежать ошибок, когда речь идёт о небелой коже. Недавнее исследование количественно оценило это для трёх коммерческих систем: ModelDerm, DeepDerm, HAM 10000. Ни одна из систем не работала лучше, чем врач-специалист, и все они продемонстрировали значительное снижение производительности между светлой и тёмной кожей. Для двух систем падение чувствительности составило около 50% в двух наборах задач (0,41→0,12, 0,45→0,25, 0,69→0,23, 0,71→0,31). Третья модель на самом деле продемонстрировала худшую чувствительность для более светлой кожи, но она также оказалась полностью несостоятельной в рабочей точке, которую использовал производитель, достигнув чувствительности <0,10 по всем направлениям. Кроме того, было обнаружено, что дерматологи, которые обычно предоставляют визуальные метки для обучения ИИ и тестовых наборов данных, также хуже работают с изображениями тёмных тонов кожи и необычных заболеваний по сравнению с аннотациями биопсии [8].

Диагностика рака молочной железы посредством маммографии, вероятно, является наиболее изученным применением компьютеров в медицинской визуализации, насчитывающей целые 10-летия. В середине 2010-х годов было выпущено несколько пакетов программного обеспечения для автоматизированного обнаружения в маммографии Computer Aided Detection (CAD), имевшее многочисленные недостатки из-за чего врачи-рентгенологи

до сих пор упускают из виду 16% случаев рака молочной железы. Данное обстоятельство могло бы послужить идеальным приложением способностей ИИ, но, несмотря на интенсивные усилия в этом направлении, продолжающиеся более 20 с лишним лет, истинный уровень рентгенолога-эксперта пока не достигнут. Многообещающие результаты в небольших исследованиях не воспроизводятся в более крупных исследованиях. Один из последних отзывов, опубликованный в сентябре 2021 года, сообщил, что 34 из 36 (94%) систем ИИ оказались менее точны, чем заключения одного врача; и все модели фреймворков были менее точными, чем консенсус двух или более специалистов. Системам искусственного интеллекта не хватило специфичности, чтобы заменить двойное чтение снимков врачом-рентгенологом в программах скрининга [9]. Стоит ли говорить, что все недостатки ИИ в радиологии, возникающие как по вине государственных, так и частных структур, подрывают доверие медицинского сообщества к ИИ. А ведь для восстановления утраченного доверия может потребоваться очень много времени и усилий.

Модель Epic Sepsis Model (ESM) была внедрена в сотнях клиник США для наблюдения за пациентами и отправки предупреждений, если они подвергались высокому риску сепсиса. Модель использует комбинацию данных мониторинга отделения неотложной помощи в режиме реального времени: частота сердечных сокращений, артериальное давление и т. д., а также демографическую информацию и информацию из медицинских карт пациента. Всего используется более 60 функций. Диагноз сепсиса устанавливался моделью на основе критериев из рекомендаций Центра по контролю и профилактике заболеваний и Международной классификации 10-го пересмотра, при наличии двух критериев, характерных для синдрома системного воспаления и одного критерия органной дисфункции, зафиксированных в течение 6 ч. Дискриминация модели оценивалась с использованием площади под кривой рабочей характеристики приёмника данных на уровне госпитализации и с горизонтами прогнозирования 4, 8, 12 и 24 ч. Внешняя проверка показала очень низкую производительность модели (AUC 0,63 против заявленных AUC 0,73 и 0,83). Из 2552 пациентов с сепсисом было выявлено только 33%, что вызвало много ложных тревог. Инициативная группа судебно-медицинских экспертов провела расследование подхода ESM к прогнозированию сепсиса, чтобы показать, как сдвиги в распределении данных могут привести к ошибкам модели машинного обучения. Было обнаружено, что изменения в стандартах кодирования способствовали снижению производительности ESM с течением времени, и что ложная корреляция в обучающих данных модели также сыграла негативную роль [10].

Самым резонансным провалом отличилась IBM со своим подразделением Watson Health. Поднятая в прессе шумиха была настолько велика, что начинали ходить слухи о «зиме ИИ». Развивая успех своей системы Watson для игры в Jeopardy, около 10 лет назад IBM запустила

Watson Health, чтобы революционизировать здравоохранение с помощью ИИ. Всё началось с разрекламированного партнёрства с Memorial Sloan Kettering для обучения ИИ на данных EHR для выработки рекомендаций по лечению. Генеральный директор IBM Дж. Рометти назвала это «нашим лунным выстрелом». На пике своего развития в Watson Health работало 7 тыс. человек. Однако совсем недавно IBM распродала всю Watson Health по частям примерно за 1 млрд долларов. Для сравнения — затраты IBM при создании Watson Health составили более 5 млрд долларов. Руководители IBM, должно быть, посчитали, что у подразделения нет абсолютно никаких шансов стать безубыточным, и решили срочно его ликвидировать. Бескомпромиссное разоблачение истории краха Watson Health можно найти на портале <https://slate.com>. Центральная тема публикации звучит так: «когда вы пытаетесь объединить браваду высоких технологий с готовностью достичь заявленных целей в сфере здравоохранения, вам предстоит предоставить абсолютно неопровержимые доказательства того, что вы в силах достичь того, о чём говорите». Предполагалось, что Watson Health изменит здравоохранение во многих важных аспектах, предоставляя онкологам информацию о лечении больных раком, фармацевтическим компаниям — о разработке лекарств, помогая проводить клинические испытания и т. п. Это звучало революционно, но на самом деле никогда не работало, потому что IBM постоянно нуждалось в огромном количестве данных для обучения модели, которых просто было невозможно найти даже за большие деньги. Партнёры IBM, такие как Онкологический центр Андерсона в Техасе, один за другим отказывались от сотрудничества, после того, как участвующие в проекте врачи жаловались, что у программы недостаёт данных, чтобы выдавать нужные рекомендации [11].

## НЕЙРОСЕТИ И ИХ ПРЕДУБЕЖДЕНИЯ

Среди многих опасений по поводу ИИ, которые привлекают внимание медицинской общественности, наиболее спорной и вместе с тем актуальной выступает проблема выявления предубеждений в алгоритмах ИИ. Ранее мы уже вскользь упоминали о неприятностях, которые связаны с проблемой систематических ошибок, которые вызваны отсутствием верифицированных датасетов, не идентифицированными алгоритмами, неправильной классификацией, погрешностями наблюдений и неграмотным обслуживанием программного обеспечения. Обоснованное беспокойство по поводу возникновения предубеждений у ИИ в ходе эксплуатации мотивировало стремление разработчиков строить справедливые свободные от ангажированности модели, что является весьма похвальным, но на деле оказывается отнюдь не лёгким делом. Справедливая модель ИИ представляется свободной от предвзятости прогностической адаптивной моделью. Она не должна быть как бы заблокированной от внешнего

мира, т. е. не должна «вариться в собственном соку». Напротив, нейросеть должна продолжать обучение, постоянно улучшая свою производительность, что со временем приведет её к реализации в качестве полноправного администратора электронной медицинской карты (ЭМК). Таким образом, будущая гипотетическая справедливая модель сможет самостоятельно функционировать в режиме поддержки принятия решений, который, однако, не будет автономным, т. е. врачи и пациенты сохранят за собой право принятия окончательного решения [12, 13].

Вместе с тем даже такая, казалось бы, максимально справедливая модель может прямо или косвенно обладать, так называемыми скрытыми предубеждениями. Точно так же, как скрытые ошибки обычно описываются как ожидающие своего проявления ошибки, в сложных программных фреймворках под *скрытыми предубеждениями* понимают ожидающие своего проявления предубеждения. Принято выделять три основные проблемы, связанные с предвзятостью в алгоритмах ИИ.

1. Первая серьёзная проблема, связанная с предвзятостью, для этого гипотетически справедливого алгоритма заключается в том, что, будучи адаптивной моделью, он со временем может стать предвзятым. Это может произойти несколькими способами. Алгоритм ИИ, обученный работать справедливо в одном контексте, способен извлечь уроки из различий в производственной деятельности медицинской организации и начать генерировать предвзятые результаты. Также нейросеть может сама учиться на широко распространённых, традиционных и порой нелепых предубеждениях, встречающихся в сфере российского здравоохранения, которые так или иначе ведут к нежелательным и даже весьма неприятным последствиям.

Скажем, алгоритм для прогнозирования смертности пациентов или индивидуальной реакции пациента на определённые методы лечения вполне может сориентироваться на существующие этнические или социально-экономические различия условий жизни каких-то групп пациентов и предсказать для них худшие результаты лечения. Но таким способом может быть создана петля отрицательной обратной связи, в результате которой предубеждения со временем будут усиливаться, что ещё больше усугубит отклоняющийся от нормы прогноз модели. С клинической точки зрения такая девиация нежелательна, т. к. корректный прогноз позволил бы перенаправить ресурсы здравоохранения именно в узкие места и дать правильные рекомендации по последующему медицинскому обслуживанию, например, по укреплению системы паллиативной помощи для социально незащищённых слоёв населения. Что ещё более важно, теперь хорошо известно, что генерация предвзятостей вполне возможна, даже если ИИ запрещено делать выводы на основе некой переменной, допустим национальности или адреса пациента, когда дата-сет не включает таковую переменную. К сожалению, это может произойти, если другие переменные коррелируют или являются

прокси для запретной переменной, что делает стратегию по исключению вызывающих озабоченность переменных бесперспективной.

2. Следующий набор проблем, связанных с предубеждениями, возникает из-за взаимодействия ИИ с клинической средой, которая включает свои собственные неявные и явные предубеждения. Следует отметить два феномена в рамках взаимодействия пациента и врача. Одним из них является феномен предвзятости автоматизации, иначе говоря, некритичное отношение к рекомендациям ИИ, которым следует неукоснительно подчиняться. Даже алгоритм, который используется просто как инструмент поддержки принятия решений, может де-факто стать настоящим автократом, если его рекомендации начнут выполняться безоговорочно. Загруженные работой и ограниченные временными рамками врачи, которые к тому же избегают юридической ответственности за игнорирование рекомендаций алгоритма, могут не обратить внимания на предвзятость ИИ. Кроме того, выделяют феномен предвзятости привилегий, т. е. непропорциональное преимущество лиц, которые уже имеют привилегии. Даже максимально честный алгоритм может быть несправедливым, если он применяется только в определённых условиях, например, в частных клиниках, обслуживающих главным образом состоятельных граждан. Классовое недоверие прекогнито к элитарным медицинским организациям, в которых в первую очередь и будет использоваться ИИ, в итоге может вылиться в повальное недоверие пациентов к его рекомендациям [14, 15].

3. Третий вид предвзятости, возможный даже у честных алгоритмов, связан с выбором цели создания модели, её заинтересованностью в определённом результате. Этот вид в чём-то напоминает первый. Однако, если интересующие исследователей исходы или выбранные для решения с помощью ИИ проблемы не отражают интересы отдельных пациентов или сообщества, это, по сути, и есть предвзятость — предпочтительный выбор или поощрение одного исхода по сравнению с другими. Иллюстрацией сказанному будет следующее. Одна из причин, по которой многие клинические испытания не смогли улучшить качество медицинской помощи, заключается в выборе суррогатов итогов исследований, которые напрямую не связаны с фактом выздоровления больного. Например, исходы лечения сердечной недостаточности оценивались только по изменениям физиологических параметров (фракции выброса левого желудочка), а не по убыванию симптоматики заболевания: снижению утомляемости и повышению толерантности к физической нагрузке. Результаты, представляющие интерес для одних групп стейкхолдеров, могут быть совсем не интересны другим и наоборот. Пациенты больше всего беспокоятся о восстановлении собственного здоровья и снижении затрат на лечение, их мало заботит эффективность системы здравоохранения как таковой. Поэтому перед внедрением алгоритмов ИИ в повседневную клиническую практику мы рекомендуем

проводить мероприятия по учёту рисков и профилактике предубеждений.

Решения ИИ с высокими рисками, например, о химиотерапевтическом лечении или проведении искусственной вентиляции лёгких, решения о материальных ресурсах, которые пациенты хотят получить от государства (социальная поддержка при инвалидности), а также автоматизированные и трудно оспариваемые решения заслуживают особо пристального внимания медицинского сообщества. Сегодня нередко встречаются модели, которые продуцируют утверждения типа: «Подобные вам, гражданин N, пациенты в аналогичной ситуации выбрали то-то и то-то». Подобный алгоритм должен считаться предвзятым из-за адаптивных предпочтений, потому что на его выбор, очевидно, могут влиять устаревшие или ненадлежащим образом интерпретированные варианты решений. К сожалению, нам не известны факты учёта опасений в отношении предвзятости ИИ. А ведь возникающие с течением времени предубеждения алгоритмов следует рассматривать как неблагоприятные события; на практике они означают, что некоторым пациентам может быть нанесён ущерб. Разрозненные действия ИИ, возникающие под влиянием предубеждений и причиняющие вред пациентам, должны стать предметом обязательной регистрации в отчётности о работе интеллектуальных медицинских устройств. Поскольку мы привыкли, что врач всегда контролирует, когда лекарственный препарат полезен одним пациентам, но вреден другим, то ожидаемо, что аналогичное требование должно быть предъявлено алгоритмам ИИ.

Аберрации в алгоритмах ИИ могут возникать не только из-за необъективности обучающих данных, но и из-за того, как нейросети обучаются с течением времени и используются на практике. Учитывая распространённость предубеждений, нет оправдания беспечному отношению к ним. Неспособность активно избавляться от предубеждений, особенно скрытых, которые проявляются неожиданно, лишь усугубляет неравенство различных групп пациентов, подрывает доверие общества к системе здравоохранения и, как это ни парадоксально, в конечном итоге, мешает ускоренному внедрению медицинского ИИ [16].

## **ДАННЫЕ АНАЛИТИЧЕСКИХ ОТЧЁТОВ МЕЖДУНАРОДНЫХ ОРГАНИЗАЦИЙ, ЗАТРАГИВАЮЩИХ ТЕМУ ТРУДНОСТЕЙ, СВЯЗАННЫХ С ТЕХНОЛОГИЯМИ МЕДИЦИНСКОГО ИСКУССТВЕННОГО ИНТЕЛЛЕКТА**

На фоне повышенного интереса к цифровому здравоохранению наблюдается большое количество внедрений ИИ без тщательного изучения доказательной базы преимуществ и вреда. Избыточный энтузиазм в отношении цифровизации привёл к распространению недолговечных

реализаций и огромному разнообразию цифровых инструментов с ограниченным пониманием их влияния на систему здравоохранения и благополучие людей. Эксперты Всемирной организации здравоохранения по этому поводу заявили: «Для улучшения здоровья и уменьшения неравенства в отношении здоровья необходима тщательная оценка электронного здравоохранения для получения фактических данных и содействия надлежащей интеграции и использованию технологий» [17].

В последнее время вышло сразу несколько интересных аналитических отчётов, затрагивающих тему искусственного интеллекта в здравоохранении. Приведём некоторые выдержки из них. KLAS Research и Колледж руководителей по управлению информацией в области здравоохранения (CHIME) опубликовали отчёт «Healthcare AI 2019. Actualizing the potential of artificial intelligence» о первых реальных кейсах встраивания систем ИИ в практическую медицину, которые касались прогнозирования повторных госпитализаций и снижения необоснованных вызовов неотложной помощи. KLAS и CHIME провели опрос среди 57 медицинских организаций, которые недавно внедрили системы на основе машинного обучения (ML) и обработки естественного языка (NLP), с целью оценки достижений в клинической, финансовой и операционной областях. KLAS оценил удовлетворённость клиентов для шести ведущих поставщиков ИИ для здравоохранения: Jvion, DataRobot, KenSci, Clinithink, IBM Watson Health and Health Catalyst. Среди прочего особое внимание в отчёте было сосредоточено на неудачах IBM Watson, авторы исследования пришли к выводу, что корпорации IBM так и не удалось исправить ситуацию со своим продуктом [18].

OptumIQ опубликовали отчёт «Annual Survey on AI in Health Care», в котором проанализировали опрос 500 руководителей организаций здравоохранения и пришли к выводу, что число внедрений ИИ в медицине увеличилось почти на 88% по сравнению с предыдущим годом. Авторы отмечают скепсис отдельных авторитетных организаторов здравоохранения в вопросе о дальнейшем росте инвестиций в ИИ, поскольку они совсем не уверены, что произведённые затраты окупятся хотя бы в течение трёхлетнего периода [19]. CB Insights опубликовала отчёт, согласно которому, несмотря на резкий рост интереса инвесторов к ИИ в здравоохранении в 2019 году, в дальнейшем возможно охлаждение инвестиционного климата. Тревожный звонок поступил от компании Freenome, использующей нейросеть для раннего выявления рака, которая в июле 2019 г. закрыла раунд финансирования на этапе В стоимостью 160 млн долларов [20].

Центр инноваций в области здравоохранения Американской ассоциации больниц выпустил отчёт «AI and Care Delivery: Emerging opportunities for artificial intelligence to transform how care is delivered». В нём исследуется использование ИИ в качестве инструмента поддержки принятия клинических решений, основываясь на мнениях экспертов в сфере здравоохранения. В докладе, в частности, рассматриваются способы решения многочисленных

проблемных вопросов, в т. ч. снижения гигантских затрат на ИИ в ходе всего цикла оказания медицинской помощи [39]. MIT Technology Review выпустил отчёт «Эффект ИИ. Как искусственный интеллект делает здравоохранение более человечным». В нём представлены данные опроса более 900 медицинских работников, проведённого MIT Technology Review Insights совместно с GE Healthcare. Обзор показал, что лишь 72% респондентов проявили прямой интерес к внедрению ИИ. 20% респондентов считают, что ИИ не сможет улучшить их экономическое положение, а 19% что ИИ не способен сделать медицинское учреждение более конкурентоспособным и клиентоориентированным. Поскольку инвестиции в медицинский ИИ заметно набирают обороты, респонденты, уже реализующие проекты глубокого обучения, с тревогой думают о том, что с каждым годом они будут тратить всё больше и больше средств на обслуживание алгоритмов. Эти выводы являются очень важны для отрасли, т. к. оказание медицинской помощи и управление ею становятся всё более сложными и дорогостоящими, а профессиональный и технологический потенциал становится всё более обременительным. На фоне того, что медики увязли в рутине постоянно увеличивающейся рабочей нагрузки и бестолковой низкооплачиваемой работы, частичная их замена чат-ботами окончательно лишит пациентов живого взаимодействия с врачами [21].

Компания KPMG выпустила отчёт «Healthcare insiders: Taking the temperature of artificial intelligence in healthcare». В нём подтверждается рост интереса к применению ИИ в медицине. Однако негативным моментом является то, что 32% респондентов не видят перспективы ИИ для объективной оценки состояния пациентов. Основными барьерами на пути ИИ названы нехватка квалифицированных кадров, высокие затраты на создание ИИ-систем и высокие риски нарушения конфиденциальности [22].

## ПРОБЛЕМНЫЕ АСПЕКТЫ ИСПОЛЬЗОВАНИЯ РАЗЛИЧНЫХ АЛГОРИТМОВ МАШИННОГО ОБУЧЕНИЯ

K-Nearest Neighbours (KNN), Random Forest (RF) и eXtreme Gradient Boosting (XGBoost) считаются самыми популярными алгоритмами машинного обучения. Они различаются своими подходами, сильными и слабыми сторонами, вариантами использования. Если быть краткими, то различие этих алгоритмов состоит в следующем.

1. KNN простой и универсальный алгоритм, используемый как для задач классификации, так и для задач регрессии. Он работает по принципу поиска  $k$ -ближайших точек данных к заданной точке запроса на основе метрики расстояния. При классификации KNN назначает класс большинства среди  $k$ -ближайших соседей в качестве прогнозируемого класса для точки запроса. В регрессии

KNN принимает среднее или средневзвешенное значение целевых значений  $k$ -ближайших соседей в качестве прогнозируемого значения для точки запроса. KNN непараметричен, т. е. не делает предположений о базовом распределении данных. Это требует больших вычислительных мощностей, особенно для больших наборов данных, ибо нужно делать расчёты расстояний для всех точек данных.

2. Random Forest — метод ансамблевого обучения, основанный на деревьях решений и в основном используемый для задач классификации и регрессии. Во время обучения он создает несколько деревьев решений, где каждое дерево обучается на случайном подмножестве функций и данных начальных выборок. При классификации окончательный прогноз основывается на большинстве голосов отдельных деревьев. Для регрессии требуется среднее предсказание отдельных деревьев. RF смягчает переобучение и обеспечивает хорошее обобщение за счёт объединения прогнозов из нескольких деревьев. Он хорошо обрабатывает многомерные данные и менее подвержен выбросам.

3. XGBoost — это усовершенствованный алгоритм повышения градиента, используемый для задач классификации, регрессии и ранжирования. Как RF, он также работает с ансамблем деревьев решений, но строит деревья последовательно, а не независимо. XGBoost использует систему повышения градиента для оптимизации ансамбля путём минимизации функции потерь. Он использует методы регуляризации, чтобы избежать переобучения и улучшить производительность модели. XGBoost эффективен в вычислительном отношении и может эффективно обрабатывать большие наборы данных. Он часто превосходит другие алгоритмы в различных соревнованиях по машинному обучению и реальных приложениях.

Применение перечисленных моделей машинного обучения зависит от конкретной задачи, для которой собираются их использовать, например, регрессии или классификации. Если говорить в целом, то нужно учитывать следующие особенности данных моделей.

K-Nearest Neighbours прост и интуитивно понятен, применим как для задач классификации, так и для задач регрессии. В KNN прогноз для новой точки данных основан на классе большинства (для классификации) или среднем значении её  $K$ -ближайших соседей (для регрессии) в пространстве признаков. Значение  $K$  — это гиперпараметр, который определяет, сколько соседних точек следует учитывать. Преимущества алгоритма:

- Легко понять и реализовать.
- Непараметрический, т. е. не делается никаких предположений относительно основного распределения данных.
- Хорошо работает на небольших наборах данных с простыми границами принятия решений.

Недостатки:

- Может оказаться дорогостоящим в вычислительном отношении для больших наборов данных, поскольку требует расчёта расстояний до всех точек данных.

- Чувствителен к несущественным функциям и шуму.
- Плохо обрабатывает несбалансированные наборы данных.

RF является методом ансамблевого обучения, который объединяет несколько деревьев решений для получения более точных прогнозов. Во время обучения он строит несколько деревьев решений и усредняет их прогнозы для повышения надёжности и точности. Каждое дерево обучается на случайном подмножестве данных и случайном подмножестве признаков, что смягчает переобучение и увеличивает обобщение. Преимущества:

- Устойчив к переобучению и хорошо работает с широким диапазоном типов данных.
- Хорошо обрабатывает многомерные данные.
- Может предоставить рейтинг важности функций.

Недостатки:

- Может быть медленным в обучении и прогнозировании больших наборов данных.
- Не хватает прозрачности и интерпретируемости по сравнению с отдельными деревьями решений.
- При выполнении некоторых сложных задач может работать не так хорошо, как более продвинутые модели типа XGBoost.

XGBoost — это расширенная реализация повышения градиента, которая представляет собой ансамблевую технику, объединяющую слабые деревья решений для создания сильной прогнозирующей модели. XGBoost совершенствует традиционное повышение градиента за счёт включения условий регуляризации, параллельной обработки и эффективной манипуляции данными для достижения более высокой точности и скорости. Преимущества:

- Высокая прогностическая эффективность благодаря механизму бустирования.
- Хорошо обрабатывает мало репрезентативные данные.
- Поддерживает регуляризацию для предотвращения переобучения.
- Быстрый и масштабируемый благодаря распараллеливанию.

Недостатки:

- Требуется настройки гиперпараметров, что может занять много времени.
- Более сложные, чем базовые модели, такие как KNN и Random Forest.
- Склонен к переобучению, если не был хорошо настроен.

Итак, подведём итоги. KNN — простой и интерпретируемый алгоритм, подходящий для небольших наборов данных, а Random Forest — мощный ансамблевый метод, обеспечивающий надёжную производительность и важность функций. XGBoost представляет собой усовершенствованный алгоритм повышения точности, который отличается высокой точностью и подходит для крупномасштабных наборов данных. Выбор модели зависит от конкретных характеристик входных данных, размера дата-сетов и желаемого баланса между простотой и производительностью прогнозирования [23, 24].

## ЗАКЛЮЧЕНИЕ

Хотя сотни алгоритмов ИИ получают одобрение от государственных надзорных органов здравоохранения в разных странах, например, в США таким органом является Управление по контролю за продуктами и лекарствами US Food and Drugs Administration (FDA), как было показано в нашем обзоре, нейросетевые платформы так или иначе проявляют склонность к скрытой предвзятости и выдаче противоречивых обобщений, особенно при недостаточности или некорректности анализируемых данных. Сохраняется слабая надежда на то, что генеративный ИИ мог бы снизить потребность в реальных данных, но его полезность, тем не менее, остаётся не до конца очевидной. Дерматологические заболевания служат весьма показательным примером создания синтетических изображений из-за разнообразия патологических проявлений, особенно с учётом цвета и тона кожи пациента. Масштабируемые алгоритмы скрытой диффузии могут генерировать изображения кожных заболеваний для дополнительного обучения модели, что, безусловно, может повысить её производительность в условиях ограниченных данных. Однако прирост производительности достигается при соотношении синтетического и реального изображений более 10:1; он существенно меньше, чем прирост, получаемый от добавления реальных изображений, поэтому сбор объективных данных остаётся главным условием для обеспечения надёжности медицинского ИИ.

Медицинский искусственный интеллект — потенциально мощный инструмент, функционирование которого сопряжено с множеством проблем. Чтобы грамотно и успешно использовать эту прогрессивную, но пока ещё несовершенную технологию, не выпуская «джина из бутылки», нам нужны эффективные стратегии и продуманное управление. Это потребует подготовки врачебных кадров совершенно нового уровня, которые смогут активно и методично участвовать в разработке, тестировании и использовании сложнейших инновационных моделей нейронных сетей. Это, в свою очередь, потребует коренного пересмотра и полного обновления программ для обучения и сертификации специалистов в области цифровой медицины, в т. ч. подрастающего поколения будущих специалистов в области цифрового здравоохранения, которые смогут гарантированно обеспечить безопасность ИИ в клинической среде. Такие шаги будут необходимы для поддержания общественного доверия к медицине в грядущую эпоху ИИ.

Несмотря на все описанные нами в статье проблемы, ясно, что ИИ станет важной частью будущего здравоохранения. Поскольку население продолжает стареть, а спрос на медицинские услуги растёт, ожидается, что нейросети будут играть решающую роль в здравоохранении — прежде всего в таких направлениях, как анализ медицинских изображений, работа виртуальных помощников,

разработка новых лекарств, формирование рекомендаций по лечению и обработка данных пациента. Усовершенствованные алгоритмы ИИ смогут анализировать снимки КТ, МРТ, ПЭТ-КТ с уровнем точности, сравнимым или даже превышающим точность специалистов-рентгенологов. Всё это в целом может помочь врачам точнее диагностировать заболевания и быстрее оценивать состояния больных, что приведёт к повышению качества и доступности оказания медицинской помощи в стране. Однако для достижения столь амбициозных целей нам потребуется реально прагматичные действия и ответственное отношение к технологиям ИИ. Создаваемая нормативно-правовая база, регулирующая алгоритмы искусственного интеллекта и машинного обучения, должна прямо включать ссылку на отслеживание отклонений в производительности, в т. ч. возникающих в процессе эксплуатации.

В заключение хотелось бы подчеркнуть, что, признавая инновационную роль, которую цифровые технологии и ИИ могут и должны играть в укреплении системы отечественного здравоохранения, нельзя упускать из виду, как важно вовремя и правильно оценивать их содействующее или негативное влияние на отрасль, чтобы обеспечивать такие управленческие решения, которые бы не отвлекали ненадлежащим образом ресурсы от альтернативных, нецифровых подходов.

## ДОПОЛНИТЕЛЬНАЯ ИНФОРМАЦИЯ

## СПИСОК ЛИТЕРАТУРЫ

1. Решетников А.В., Шамшурина Н.Г., Шамшурин В.И. Экономика и управление в здравоохранении. 2-е изд. Москва: Издательство Юрайт, 2020. EDN: KSZBPT
2. Reshetnikov A., Fedorova J., Prisyazhnaya N., et al. Health management for sustainable development. В кн.: 2018 Second World Conference on Smart Trends in Systems, Security and Sustainability (WorldS4). IEEE, 2018.
3. Berdutin V. Socionic vision on Bioethics and Deontology. Lap Lambert Academic Publishing, 2018.
4. Liu J. Artificial Intelligence and Data Analytics Applications in Healthcare General. Review and Case Studies. In: CAIH2020: Proceedings of the 2020 Conference on Artificial Intelligence and Healthcare; Oct 2020. P. 49–53. doi: 10.1145/3433996.3434006
5. Daley K. Two arguments against human-friendly AI // AI and Ethics. 2021. Vol. 1, N 4. P. 435–444. doi: 10.1007/s43681-021-00051-6
6. Vellido A. Societal Issues Concerning the Application of Artificial Intelligence in Medicine // Kidney Dis. 2019. Vol. 5, N 1. P. 11–17. doi: 10.1159/000492428
7. Breede E., Bayor E., Hersh F., et al. A Human-Centered Evaluation of a Deep Learning System Deployed in Clinics for the Detection of Diabetic Retinopathy. In: CHI 2020; 2020 Apr 25–30; Honolulu. P. 1–12. doi: 10.1145/3313831.3376718
8. Daneshjou R., Vodrahalli K., Novoa R.A., et al. Disparities in Dermatology AI Performance on a Diverse, Curated Clinical Image Set [Internet]. Cornell University, 2022. Режим доступа: <https://arxiv.org/ftp/arxiv/papers/2203/2203.08807.pdf> Дата обращения: 05.09.2023.

Настоящая статья представляет собой продолжение статьи: Бердутин В.А., Романова Т.Е., Романов С.В., Абаева О.П. Проблемы медицинского искусственного интеллекта. Часть 1 // Социология медицины. 2023. Т. 22, № 2. С. 202–211. DOI: <https://doi.org/10.17816/socm619132>

**Конфликт интересов.** Авторы заявляют об отсутствии конфликта интересов.

**Источник финансирования.** Отсутствует.

**Вклад авторов.** Все авторы подтверждают соответствие своего авторства международным критериям ICMJE. Наибольший вклад распределён следующим образом: все авторы — концепция обзора, сбор и обработка материала; В.А. Бердутин — написание текста; Т.Е. Романова — редактирование.

## ADDITIONAL INFORMATION

This article is a continuation of the article by Berdutin VA, Romanova TE, Romanov SV, Abaeva OP. Problematic aspects of medical artificial intelligence. Part 1. Sociology of Medicine. 2023;22(2):202–211. DOI: <https://doi.org/10.17816/socm619132>

**Competing interests.** The authors declare that they have no competing interests.

**Funding source.** Not specified.

**Author's contribution.** All authors confirm compliance of their authorship with the international ICMJE criteria. The largest contribution is distributed as follows: all authors — review concept, collection and processing of materials; V.A. Berdutin — writing the text; T.E. Romanova — editing.

[arxiv.org/ftp/arxiv/papers/2203/2203.08807.pdf](https://arxiv.org/ftp/arxiv/papers/2203/2203.08807.pdf) Дата обращения: 05.09.2023.

9. Freeman K., Geppert J., Stinton Ch., Todkill D., et al. Use of artificial intelligence for image analysis in breast cancer screening programs: systematic review of test accuracy // BMJ. 2021. Vol. 374. P. n1872. doi: 10.1136/bmj.n1872
10. Wong A., Otles E., Donnelly J.P., et al. External Validation of a Widely Implemented Sepsis Prediction Model in Hospitalized Patients // JAMA Intern Med. 2021. Vol. 181, N 8. P. 1065–1070. doi: 10.1001/jamainternmed.2021.2626
11. O'Leary L. How IBM's Watson Went from the Future of Health Care to Sold Off for Parts. B: Slate [интернет]. 2022. Режим доступа: <https://slate.com/technology/2022/01/ibm-watson-health-failure-artificial-intelligence.html> Дата обращения: 23.09.2023
12. Khan B., Hajira F., Qureshi A., et al. Drawbacks of Artificial Intelligence and Their Potential Solutions in the Healthcare Sector. In: Biomedical Materials & Devices, 2023. Feb 8. P. 1–8. doi: 10.1007/s44174-023-00063-2
13. Lee T.T., Kesselheim A.S. U.S. Food and Drug Administration Precertification Pilot Program for Digital Health Software: Weighing the Benefits and Risks // Ann Intern Med. 2018. Vol. 168, N 10. P. 730–732. doi: 10.7326/M17-2715
14. Parikh R.B., Teeple S., Navathe A.S. Addressing bias in artificial intelligence in health care // JAMA. 2019. Vol. 322, N 24. P. 2377–2378. doi: 10.1001/jama.2019.18058

15. Challen R., Denny J., Pitt M., et al. Artificial intelligence, bias and clinical safety // *BMJ Qual Saf.* 2019. Vol. 28, N 3. P. 231–237. doi: 10.1136/bmjqs-2018-008370
16. He J., Baxter S.L., Xu J., et al. The practical implementation of artificial intelligence technologies in medicine // *Nat. Med.* 2019. Vol. 25, N 1. P. 30–36. doi: 10.1038/s41591-018-0307-0
17. Monitoring the implementation of digital health: an overview of selected national and international methodologies [Internet]. Copenhagen: WHO Regional Office for Europe, 2022. Режим доступа: <https://www.who.int/europe/publications/i/item/WHO-EURO-2022-5985-45750-65816> Дата обращения: 20.09.2023.
18. Gale A. Reimagined Hospitals. How Far Is the Future? // *HealthManagement.org The Journal.* 2020. Vol. 20, N 1. P. 36–38.
19. Christensen J. A Snapshot of Imaging Technology: Exciting Developments and When to Expect Them // *HealthManagement.org The Journal.* 2020. Vol. 20, N 6. P. 476–479
20. Landi H. Investors poured \$4B into healthcare AI startups in 2019. B: *Fierce Healthcare* [интернет]. Questex, 2020. Режим дос-

- тупа: <https://www.fiercehealthcare.com/tech/investors-poured-4b-into-healthcare-ai-startups-2019> Дата обращения: 23.09.2023
21. Memora Health raises \$40M for its virtual care delivery platform. Memora Health competitors include Wheel, Welby Health, and Twistle. ResearchBriefs. B: *CBinsights* [интернет]. 2022. Режим доступа: <https://www.cbinsights.com/research/мемора-health-competitors-wheel-welby-health-twistle/> Дата обращения: 24.09.2023
22. The AI effect: How artificial intelligence is making health care more human. B: *Technology review* [интернет]. GE Healthcare. Режим доступа: <https://www.technologyreview.com/hub/ai-effect/> Дата обращения: 13.09.2023
23. Avuçlu E. Determining the most accurate machine learning algorithms for medical diagnosis using the monk' problems database and statistical measurements. *Journal of Experimental & Theoretical Artificial Intelligence.* Forthcoming. 2023. doi: 10.1080/0952813X.2023.2196984
24. Shukla S. Enhancing healthcare insights, exploring diverse use-cases with K-means clustering // *International Journal of Management, IT & Engineering.* 2023. Vol. 13, N 8. P. 60–68.

## REFERENCES

1. Reshetnikov AV, Shamshurina NG, Shamshurin VI. *Economics and management in healthcare: Textbook and workshop.* 2<sup>nd</sup> ed. Moscow: Yurayt Publishing House; 2020 (In Russ.) EDN: KSZBPT
2. Reshetnikov A, Fedorova J, Prisyazhnaya N, et al. Health management for sustainable development. In: *2018 Second World Conference on Smart Trends in Systems, Security and Sustainability (WorldS4).* IEEE, 2018.
3. Berdutin V. *Socionic vision on Bioethics and Deontology.* LAP LAMBERT Academic Publishing; 2018.
4. Liu J. Artificial Intelligence and Data Analytics Applications in Healthcare General. Review and Case Studies. In: *CAIH2020: Proceedings of the 2020 Conference on Artificial Intelligence and Healthcare, October 2020*, P. 49–53. doi: 10.1145/3433996.3434006
5. Daley K. Two arguments against human-friendly AI. *AI and Ethics.* 2021;1(4):435–444. doi: 10.1007/s43681-021-00051-6
6. Vellido A. Societal Issues Concerning the Application of Artificial Intelligence in Medicine. *Kidney Dis.* 2019;5(1):11–17. doi: 10.1159/000492428
7. Breede E, Bayor E, Hersh F, et al. *A Human-Centered Evaluation of a Deep Learning System Deployed in Clinics for the Detection of Diabetic Retinopathy.* In: *CHI 2020*; 2020 Apr 25–30; Honolulu. P. 1–12. doi: 10.1145/3313831.3376718
8. Daneshjou R, Vodrahalli K, Novoa RA, et al. *Disparities in Dermatology AI Performance on a Diverse, Curated Clinical Image Set* [Internet]. Cornell University; 2022. [cited 2023 Sep 05]. Available from: <https://arxiv.org/ftp/arxiv/papers/2203/2203.08807.pdf>
9. Freeman K, Geppert J, Stinton Ch, et al. Use of artificial intelligence for image analysis in breast cancer screening programs: systematic review of test accuracy. *BMJ.* 2021;374:n1872. doi: 10.1136/bmj.n1872
10. Wong A, Otles E, Donnelly JP, et al. External Validation of a Widely Implemented Sepsis Prediction Model in Hospitalized Patients. *JAMA Intern Med.* 2021;181(8):1065–1070. doi: 10.1001/jamainternmed.2021.2626
11. O'Leary L. How IBM's Watson Went from the Future of Health Care to Sold Off for Parts. In: *Slate* [Internet]. 2022. [cited 2023 Sep 23]. Available from: <https://slate.com/technology/2022/01/ibm-watson-health-failure-artificial-intelligence.html>
12. Khan B, Hajira F, Qureshi A, et al. Drawbacks of Artificial Intelligence and Their Potential Solutions in the Healthcare Sector. In: *Biomedical Materials & Devices*; 2023. Feb 8. P. 1–8. doi: 10.1007/s44174-023-00063-2
13. Lee TT, Kesselheim AS. U.S. Food and Drug Administration Precertification Pilot Program for Digital Health Software: Weighing the Benefits and Risks. *Ann Intern Med.* 2018;168(10):730–732. doi: 10.7326/M17-2715
14. Parikh RB, Teeple S, Navathe AS. Addressing Bias in Artificial Intelligence in Health Care. *JAMA.* 2019;322(24):2377–2378. doi: 10.1001/jama.2019.18058
15. Challen R, Denny J, Pitt M, et al. Artificial intelligence, bias and clinical safety. *BMJ Qual Saf.* 2019;28(3):231–237. doi: 10.1136/bmjqs-2018-008370
16. He J, Baxter SL, Xu J, et al. The practical implementation of artificial intelligence technologies in medicine. *Nat Med.* 2019;25(1):30–36. doi: 10.1038/s41591-018-0307-0
17. *Monitoring the implementation of digital health: an overview of selected national and international methodologies* [Internet]. Copenhagen: WHO Regional Office for Europe; 2022. [cited 2023 Sep 20]. Available from: <https://www.who.int/europe/publications/i/item/WHO-EURO-2022-5985-45750-65816>
18. Gale A. Reimagined Hospitals. How Far Is the Future? *HealthManagement.org The Journal.* 2020;20(1):36–38.
19. Christensen J. A Snapshot of Imaging Technology: Exciting Developments and When to Expect Them. *HealthManagement.org The Journal.* 2020;20(6):476–479.
20. Landi H. Investors poured \$4B into healthcare AI startups in 2019. In: *Fierce Healthcare* [Internet]. Questex; 2020 [cited 2023 Sep 23]. Available from: <https://www.fiercehealthcare.com/tech/investors-poured-4b-into-healthcare-ai-startups-2019>
21. Memora Health raises \$40M for its virtual care delivery platform. Memora Health competitors include Wheel, Welby Health, and Twistle. ResearchBriefs. In: *CBinsights* [Internet]; 2022. [cited 2023 Sep 24] Available from: <https://www.cbinsights.com/research/memora-health-competitors-wheel-welby-health-twistle/>

**22.** The AI effect: How artificial intelligence is making health care more human. In: *Technology review* [Internet]. GE Healthcare. [cited 2023 Sep 13]. Available from: <https://www.technologyreview.com/hub/ai-effect/>

**23.** Avuçlu E. Determining the most accurate machine learning algorithms for medical diagnosis using the monk' problems

database and statistical measurements. *Journal of Experimental & Theoretical Artificial Intelligence*. Forthcoming. 2023. doi: 10.1080/0952813X.2023.2196984

**24.** Shukla S. Enhancing Healthcare Insights, Exploring Diverse Use-Cases with K-means Clustering. *International Journal of Management, IT & Engineering*. 2023;13(8):60–68.

## ОБ АВТОРАХ

\* **Бердугин Виталий Анатольевич**, канд. мед. наук;  
адрес: Россия, 123098, Москва, ул. Живописная, д. 46;  
ORCID: 0000-0003-3211-0899;  
eLibrary SPIN: 8316-7111;  
e-mail: vberdt@gmail.com

**Романова Татьяна Евгеньевна**, канд. мед. наук;  
ORCID: 0000-0001-6328-079X;  
eLibrary SPIN: 4943-6121;  
e-mail: drmedromanova@gmail.com

**Романов Сергей Владимирович**, д-р мед. наук;  
ORCID: 0000-0002-1815-5436;  
eLibrary SPIN: 9014-6344;  
e-mail: director@pomc.ru

**Абаева Ольга Петровна**, д-р мед. наук, проф.;  
ORCID: 0000-0001-7403-7744;  
eLibrary SPIN: 5602-2435;  
e-mail: abaevaop@inbox.ru

## AUTHORS' INFO

\* **Vitalii A. Berdutin**, MD, Cand. Sci. (Medicine);  
address: 46 Zhivopisnaya street, 123098 Moscow, Russia;  
ORCID: 0000-0003-3211-0899;  
eLibrary SPIN: 8316-7111;  
e-mail: vberdt@gmail.com

**Tatyana E. Romanova**, MD, Cand. Sci. (Medicine);  
ORCID: 0000-0001-6328-079X;  
eLibrary SPIN: 4943-6121;  
e-mail: drmedromanova@gmail.com

**Sergey V. Romanov**, MD, Dr. Sci. (Medicine);  
ORCID: 0000-0002-1815-5436;  
eLibrary SPIN: 9014-6344;  
e-mail: director@pomc.ru

**Olga P. Abaeva**, MD, Dr. Sci. (Medicine), Professor;  
ORCID: 0000-0001-7403-7744;  
eLibrary SPIN: 5602-2435;  
e-mail: abaevaop@inbox.ru

\* Автор, ответственный за переписку / Corresponding author