

ЭТИКА ПРИМЕНЕНИЯ LLM-MODEЛЕЙ В МЕДИЦИНЕ И НАУКЕ

Л. Ф. Габидуллина¹ ✉, М. Ю. Котловский^{1,2}¹ Ярославский государственный медицинский университет, Ярославль, Россия² Национальный научно-исследовательский институт общественного здоровья имени Н. А. Семашко, Москва, Россия

Быстрое внедрение больших языковых моделей (LLM) в здравоохранение порождает острые этические дилеммы и практические риски. Центральная проблема связана с доверием медицинских специалистов, пациентов и разработчиков к этим системам, а также с потенциальным нарушением основополагающих принципов медицинской этики. Данное изложение анализирует ключевые вызовы, включая критическую важность доверия (зависящую от качества данных LLM), нарушение информированного согласия и автономии пациента из-за отсутствия прозрачности и чрезмерной опоры на ИИ. Особое внимание уделяется рискам защиты конфиденциальных медицинских данных, что подтверждается инцидентами несанкционированной передачи информации при использовании общедоступных LLM. Необходимость разработки прозрачных, безопасных и этически регулируемых решений для LLM в медицине становится первостепенной задачей.

Ключевые слова: этические дилеммы, большие языковые модели (LLM)

Вклад авторов: Л. Ф. Габидуллина — анализ литературы, планирование исследования, написание текста, редактирование; М. Ю. Котловский — сбор, анализ, интерпретация данных.

✉ **Для корреспонденции:** Ландыш Фаритовна Габидуллина
ул. Революционная, д. 5, г. Ярославль, 150000, Россия; landush10@yandex.ru

Статья поступила: 23.07.2025. **Статья принята к печати:** 05.09.2025 **Опубликована онлайн:** 22.09.2025

DOI: 10.24075/medet.2025.016

ETHICS OF APPLYING LLM-MODELS IN MEDICINE AND SCIENCE

Gabidullina LF¹ ✉, Kotlovsky MY^{1,2}¹ Yaroslavl State Medical University, Yaroslavl, Russia² Department of Strategic Analysis in Healthcare N. A. Semashko National Research Institute of Public Health, Moscow, Russia

Rapid integration of large language models (LLM) into healthcare gives rise to acute ethical dilemmas and practical risks. The principal issue is associated with trust of medical professionals, patients and developers in the models, as well as with the potential violation of medical ethics. In the article, key challenges are analyzed including critical importance of trust (depending on LLM data quality), disturbance of informed consent and autonomy of a patient due to the lack of transparency and excessive trust in AI algorithms. Particular attention is given to the risks of confidential medical data protection, which is confirmed by non-authorized transfer of data while using generally accessible LLM. The need to develop transparent, safe and ethically regulated solutions for LLM in medicine is prioritized.

Key words: ethical dilemmas, large language models (LLM)

Author contribution: Gabidullina LF — literature analysis, research planning, writing, editing; Kotlovsky MY — data collection, analysis, and interpretation.

✉ **Correspondence should be addressed:** Landush F Gabidullina,
Revolyutsionnaya str., 5, Yaroslavl, 150000, Russia; landush10@yandex.ru

Received: 23.07.2025 **Accepted:** 05.09.2025 **Published online:** 22.09.2025

DOI: 10.24075/medet.2025.016

Мы часто сталкиваемся с цифровыми технологиями в своей работе, обращаемся к искусственному интеллекту, но уверены ли мы что всегда действуем этично и безопасно в этом цифровом мире? Мы убеждены в том, что интеграция передовых технологий, в частности больших языковых моделей (LLM), в сферу здравоохранения и науки требует не только технических знаний, но и глубокого этического осмысления. Наша статья посвящена именно этой актуальной теме.

Одной из разновидностей искусственного интеллекта являются большие языковые модели (LLM, Large Language Models), основанная на архитектуре трансформеров, обученная на огромных массивах текстовых данных для выполнения широкого спектра задач, связанных с обработкой естественного языка (NLP), таких как генерация текста, перевод, суммаризация, вопросно-ответные системы, кодирование и другие.

Они обучаются в режиме самопредсказания (например, предсказания следующего слова). Тренируются на

триллионах токенов из различных источников (интернет, книги, научные статьи и др.). Содержат миллиарды параметров (например, GPT-3-175 млрд, GPT-4 — оценки от 500 млрд). Могут адаптироваться под задачи без дополнительного обучения (zero-shot, few-shot, fine-tuning). Выдают человекообразный текст, но не обладают «пониманием» в человеческом смысле [1].

Обучение LLM моделей проходит следующим образом:

- сбор данных;
- токенизация (Byte Pair Encoding (BPE) или SentencePiece);
- предобучение (Pretraining);
- дополнительная настройка (Fine-tuning).

Этапы работы при использовании LLM:

- формулировка запроса (Prompting);
- кодирование запроса проходит через токенизацию. Токены преобразуются в векторы — эмбединги;
- обработка моделью (Inference). Векторы проходят через слои трансформера. На каждом слое

учитывается контекст всех предыдущих токенов через self-attention. Модель генерирует вероятности для следующего токена → выбирает наиболее вероятный;

- формирование ответа;
- постобработка и фильтрация (в продакшене). Иногда подключаются дополнительные модули (ретриверы, базы знаний и т.д.).

Задача распознавания речи заключается в преобразовании устной речи, записанной на аудиофайле, в текстовую форму. Синтез речи по-прежнему остаётся сложной задачей, особенно когда необходимо добиться естественного звучания и передачи эмоций.

Задача генерации изображений по тексту состоит в том, чтобы создать изображение по его текстовому описанию. Эта задача остаётся достаточно сложной, однако в последние годы широкое распространение получили диффузионные модели — архитектуры, используемые для генерации изображений в области компьютерного зрения. Особую популярность они приобрели с 2020 года. На базе диффузионных моделей в 2021 году была создана нейросеть MidJourney, а в июне 2022 года Сбер представил первую версию собственной генеративной нейросети под названием Kandinsky.

Обе нейросети показывают очень хорошее качество генерации изображений на основе текста.

Вопросно-ответные системы (Question-Answering Systems) представляют собой модели машинного обучения, способные находить ответы на вопросы на основе заданного текста. Таким образом, LLM демонстрируют поразительные способности в обработке и генерации текста, анализе данных, а также в создании рекомендаций.

ПАЦИЕНТЫ И МЕТОДЫ

Существует более сорока различных моделей LLM, растёт из года в год их финансирование и в 2022 году они обогнали по ассоциативному мышлению средние показатели человека. Однако, с этим потенциалом приходят и новые, зачастую сложные, этические дилеммы. В контексте здравоохранения, где на кону стоит человеческое здоровье и жизнь, эти вопросы приобретают особую остроту. Одним из центральных аспектов, требующих пристального внимания, является вопрос доверия медицинских специалистов к инструментам на базе LLM, а пациентов к системе, использующей эти технологии, и даже доверия разработчиков к надёжности своих систем в критически важных условиях [2].

Обученные на колоссальных объемах данных, часто закрытых или недостаточно верифицированных, не имеющих прозрачности и полной объяснимости алгоритмических решений, не всегда могут выдавать причины своего вывода и объяснить логику. В условиях клинической практики такая непрозрачность становится источником потенциального риска, если рекомендации ассистента принимаются без критического осмысления со стороны врача. Можно ли доверять инструменту, чьи решения необъяснимы? Этическая практика требует как минимум минимального уровня прозрачности для любого инструмента, влияющего на здоровье человека.

Вопрос доверия тесно связан с качеством и репрезентативностью данных, на которых обучаются LLM. Медицинские данные могут быть не только неполными, но и содержать искажения, предвзятости или ошибки, обусловленные спецификой сбора анамнеза, субъективной интерпретацией симптомов, региональными различиями в медицинских подходах или

даже социально-экономическим статусом пациентов. Если LLM извлекает закономерности из таких некорректных или предвзятых данных, это может привести к генерации рекомендаций, не соответствующих лучшим клиническим практикам или принципам медицинской этики [3].

Например, модель может «переобобщить», предложив стандартное лечение, игнорируя индивидуальные особенности или сопутствующие заболевания пациента, что потенциально может нанести вред.

Не менее важным является вопрос восприятия LLM-ассистентов пациентами. Если пациент не проинформирован о том, что в процессе его диагностики или формирования рекомендаций по лечению использовалась машинная помощь, это является прямым нарушением принципа информированного согласия и автономии. Этический императив взаимодействия требует полной открытости: пациент должен быть осведомлен об участии искусственного интеллекта в своем лечении и иметь право отказаться от такого подхода.

Помимо клинической практики, большие языковые модели активно внедряются в медицинские научные исследования. Они используются для анализа огромных объемов научных публикаций, генерации гипотез, поддержки в дизайне экспериментов, суммаризации результатов, и даже для первичного анализа данных в рамках доклинических и клинических исследований. При расширении новых горизонтов использования LLM-ассистентов появляются новые риски, связанные с этическими вопросами, такие как:

- достоверность и «галлюцинации» (полностью вымышленные, но грамотно оформленные утверждения; некорректные рекомендации (например, лекарств с противопоказаниями); ложные источники или ссылки (не существующих научных публикаций);
- предвзятость в исследованиях;
- вопросы авторства и интеллектуальной собственности;
- воспроизводимость и верификация;
- конфиденциальность данных.

Поэтому этическое регулирование требует расширенного комплексного подхода и междисциплинарного сотрудничества, т.е.:

- мы должны стремиться к созданию LLM, способных объяснить логику своих решений;
- необходимо минимизировать предвзятость и ошибки в данных, используемых для обучения моделей;
- пациенты и участники исследований должны быть полностью информированы об использовании ИИ и иметь право выбора;
- четко определить ответственности разработчика, врача, который использовал систему, или самой системы в случае ошибки;
- создание рекомендаций по использованию LLM в науке, включая вопросы авторства, достоверности и воспроизводимости;
- медицинские специалисты и ученые должны быть обучены грамотному и критическому использованию LLM, пониманию их ограничений и этических аспектов.

Если обширные обучающие выборки содержат исторически сложившиеся предубеждения, неточности или отражают системные неравномерности, присущие реальной клинической практике, модели могут не только их усвоить, но и невольно тиражировать или даже усугублять.

Возьмем, к примеру, этническую предвзятость. Если в обучающих данных исторически недооценивались симптомы или чаще игнорировались потребности пациентов определенных этнических групп, LLM может перенять и воспроизвести эти нежелательные паттерны.

Аналогично, гендерное смещение. Долгое время медицинские исследования акцентировали внимание на мужском организме как «стандартном», игнорируя специфику женского здоровья.

Риски также распространяются на пациентов с ограниченными возможностями, психическими расстройствами или представителей уязвимых групп. Если LLM неосознанно воспроизводит стереотипы, она может стать источником непреднамеренной дискриминации в клинической практике [4].

Особая опасность кроется в скрытом характере этих смещений. Из-за проблемы непрозрачности и отсутствия полной объяснимости алгоритмических решений, о которой говорили ранее, LLM не способна объяснить логику своих решений. Ее рекомендации могут выглядеть убедительно и точно, что может создавать ложное ощущение объективности у врача. Более того, алгоритмическая предвзятость редко проявляется в единичных случаях; она становится очевидной лишь при агрегированном анализе больших массивов данных. Однако для отдельного пациента ошибочное решение может иметь необратимые последствия.

Многие законодательные акты обязывают медицинские учреждения соблюдать высочайшие стандарты в отношении хранения, обработки, доступа и передачи данных пациентов. Однако с появлением в этом процессе сторонней технологической системы, такой как LLM, точки потенциального нарушения безопасности и конфиденциальности возрастают многократно [5].

Нарушение конфиденциальности может повлечь:

- психологический и социальный вред пациенту;
- утрату доверия к системе здравоохранения;
- юридическую ответственность.

LLM, особенно при обучении и использовании в медицине, могут собирать, обрабатывать, воспроизводить или непреднамеренно раскрывать такие данные напрямую или по косвенным признакам.

Особую тревогу вызывает тот факт, что многие языковые модели, особенно те, что предлагаются коммерческими разработчиками, могут иметь встроенные механизмы логирования всех взаимодействий. Даже если утверждается, что информация «обезличена», существует высокая вероятность реидентификации пациента при проведении сложных корреляций данных. Это абсолютно недопустимо в контексте базовых этических принципов медицинской практики.

Стоит ли медицинскому специалисту жертвовать абсолютной конфиденциальностью данных пациента ради более точного анализа симптомов, выполненного на основе машинного интеллекта?

Пациенты должны не просто подписывать шаблонные формы, а получать доступную, понятную и исчерпывающую информацию о том, в каком виде, где и как их данные будут использоваться LLM-системами, какие потенциальные риски с этим связаны, и кто будет иметь к ним доступ. Этическое управление данными выходит далеко за рамки формальной юридической защиты.

Создание надежных технических и организационных протоколов, обеспечивающих абсолютную конфиденциальность и безопасность данных, должно

стать не просто желаемым условием, а неотъемлемым требованием их этического внедрения. В противном случае, даже самая точная и потенциально «полезная» модель может превратиться в источник глубокого нарушения базовых прав пациента и подорвать фундамент доверия, который мы стремимся построить.

Продолжая разговор о доверии, которое, как мы выяснили, неразрывно связано с прозрачностью и защитой данных, стоит упомянуть о влиянии LLM на автономию пациента и саму суть врачебной этики.

Информационная асимметрия и отчуждение пациента: если врач, полагаясь на авторитет ИИ, не будет тщательно переводить эту информацию на понятный пациенту язык и вовлекать его в дискуссию, мы рискуем фактически «маргинализировать» пациента в процессе принятия решений, лишая его истинного информированного согласия.

Эрозия врачебной субъектности и ответственности: традиционно врач — это не только носитель знаний, но и моральный актор, руководствующийся эмпатией, состраданием и глубокой личной ответственностью. Искусственный интеллект, несмотря на свою аналитическую мощь, лишен этих человеческих качеств; его рекомендации основаны на алгоритмах и статистике, а не на моральной оценке или понимании уникальной человеческой ситуации. Чрезмерная зависимость врача от выводов LLM, восприятие их как неоспоримого «объективного» авторитета, может привести к снижению критического мышления врача и ослаблению его этической позиции как лица, несущего окончательную ответственность за принятие решений.

Влияние на доверие пациента к врачу: это напрямую подводит нас к уже затронутой теме «доверия». Если пациент ощущает, что ключевые решения, касающиеся его здоровья, принимаются или сильно зависят от «машины», а не от живого врача, его вера в искреннюю заботу и индивидуальный подход может быть подорвана.

Ограничение выбора и стандартизация решений: LLM, оптимизированные для выдачи «оптимальных» или статистически обоснованных рекомендаций, создают риск подмены «индивидуализированного подхода» к лечению стандартизированными протоколами, что особенно опасно в условиях автоматизированного триажа или при ограниченном доступе к непосредственному врачебному контакту.

РЕЗУЛЬТАТЫ ИССЛЕДОВАНИЯ

Медицинская дезинформация в ответах ChatGPT

В исследовании (Ayers, et al., 2023) сравнили ответы врачей и ChatGPT на реальные медицинские вопросы пациентов. Хотя по тону ответы ИИ были более «вежливыми», в 27% случаев они содержали потенциально опасные или неточные рекомендации.

Реакция: предупреждение от JAMA и призывы не использовать ChatGPT без верификации в телемедицине [6].

Фейковые новости и генерация дезинформации в биомедицинской сфере

Исследователи протестировали GPT на способность создавать дезинформацию: модель легко генерировала фальшивые новости о «новом вирусе», «вакцине, вызывающей рак» и т.д.

Реакция: ООН и ВОЗ призвали к ограничению использования ИИ в области общественного здравоохранения без этической экспертизы [7].

Фейковые статьи и источники в студенческих и научных работах

В 2023–2024 гг. во многих университетах мира студенты начали массово сдавать работы, содержащие вымышленные библиографические ссылки, сгенерированные LLM.

Результат: университеты начали вводить официальные запреты на безотчетное использование LLM, усилили антиплагиатную проверку.

Университет в Мельбурне аннулировал диплом одного из выпускников после обнаружения вымышленных источников в магистерской диссертации [8].

Утечка данных пациентов через ChatGPT в Samsung (2023)

Сотрудники Samsung в Южной Корее использовали ChatGPT для обработки внутренней документации, включая медицинские записи и анализ кода диагностики в биомедицинском ПО.

Проблема: переписка с LLM сохраняется на серверах OpenAI и может быть использована для обучения моделей, если не отключены соответствующие настройки. Это создало угрозу утечки чувствительных данных.

Результат: Samsung ввела запрет на использование публичных LLM. Компания начала разработку собственной офлайн-модели [9].

NHS Великобритания: использование GPT через сторонние интерфейсы без верификации

В ряде госпиталей врачи стали использовать публичные веб-версии ChatGPT для ускорения подготовки выписок, рецептов и аннотаций. Иногда они копировали фрагменты из истории болезни в интерфейс чата.

Риски: данные пересылались на серверы за пределами юрисдикции Великобритании (и GDPR), нарушая законы о защите медицинской тайны.

Результат: NHS выпустила экстренное указание не использовать открытые LLM, пока не будут внедрены защищенные решения [10].

ОБСУЖДЕНИЕ РЕЗУЛЬТАТОВ

Чем же нужно руководствоваться при соблюдении этики применения технологий и ИИ в медицине?

1. Принципом «человек превыше всего». Все решения и разработки должны быть направлены на благо пациента, защиту его прав и достоинства, а не на максимизацию технологических возможностей ради них самих.
2. Прозрачность, открытость в отношении того, как собираются, используются и защищаются данные, а также как и в какой степени ИИ участвует в процессе диагностики и лечения, является фундаментом для информированного согласия и доверия.
3. Необходимо четко определить, кто несет окончательную ответственность за решения, принятые с участием ИИ. В медицине эта ответственность всегда должна оставаться за человеком — врачом.
4. Этика применения ИИ также требует предотвращения дискриминации и предвзятости, которые могут быть присущи алгоритмам, и обеспечения равного доступа к качественной помощи.
5. Технологии должны расширять, а не ограничивать право пациента на информированный выбор и активное участие в управлении своим здоровьем.

Предлагаем несколько ключевых направлений для дальнейшего развития.

1. Разработка национальных и международных этических стандартов и руководств для применения LLM в здравоохранении, которые будут учитывать не только технологические, но и социокультурные особенности.
2. Интеграция этики ИИ в медицинское образование и непрерывное профессиональное развитие. Будущие и практикующие врачи должны быть обучены не только работе с технологиями, но и их этическим аспектам.
3. Создание мультидисциплинарных команд (врачи, этики, юристы, инженеры, представители пациентов) для постоянного мониторинга, оценки и адаптации этических принципов к быстро меняющемуся технологическому ландшафту.
4. Приоритизация исследований, направленных на понимание долгосрочного воздействия ИИ на отношения «врач-пациент» и на психоэмоциональное состояние обеих сторон.
5. Активное внедрение практик совместного принятия решений, где ИИ предоставляет информацию, но окончательный выбор всегда остается за человеком, в диалоге с врачом.

Практические и этические рекомендации:

- запрет использования открытых LLM для ввода ПМД (персональных медицинских данных) без специальных соглашений;
- анонимизация данных перед обработкой;
- локальные или защищенные LLM-сервисы, развернутые внутри клиники;
- разработка этических протоколов согласия на обработку данных с ИИ;
- обязательная аудитируемость использования ИИ в медицинских ИС;
- шифрование и логирование всех обращений к LLM при работе с пациентами.

Подводя итог нашему обсуждению, мы можем утверждать, что внедрение LLM в здравоохранение — это не просто технологический прорыв, но и глубокий этический вызов, требующий осознанного и ответственного подхода. Как мы выяснили, краеугольным камнем здесь выступает доверие пациента к системе, к врачу и, в конечном счете, к самой технологии. Это доверие не может быть достигнуто без строжайшего соблюдения принципов, которые мы сегодня затронули.

ВЫВОДЫ

Этические аспекты применения больших языковых моделей (LLM) в медицине представляют собой сложный и многоуровневый вызов. Вопросы надежности, достоверности информации, ответственности за ошибки, конфиденциальности данных, предвзятости и дискриминации, прозрачности и объяснимости, а также неравного доступа требуют внимательного анализа и строгого регулирования. Только комплексный подход, включающий технические, юридические и социальные меры, позволит минимизировать риски и максимально эффективно использовать потенциал LLM в клинической практике.

Таким образом, наш путь в цифровую эру медицины должен быть освещен не только светом инноваций, но пламенем этических принципов, гарантирующих, что технологии служат человеку, его здоровью и благополучию, а не наоборот.

Литература

1. Косарев Е. А. Учебник по машинному обучению. Режим доступа: [Электронный ресурс] <https://education.yandex.ru/handbook/ml/article/yazykovye-modeli> (дата обращения: 24.05.2025).
2. Хохлов А. Л., Котловский М. Ю., Павлов А. В., Потапов М. П., Габидуллина Л. Ф., Цыбикова Э. Б. Развитие нейротехнологий: этические проблемы и общественные дискуссии. Медицинская этика. 2024; 1: 20–25.
3. Этические вопросы LLM-ассистентов в клинике [Интернет]. Медицинские новости и статьи. 2025, май — [по состоянию на 18 июня 2025 года]. Режим доступа: [Электронный ресурс] <https://nexusacademy.ru/tpost/95kvl3v0k1-eticheskie-voprosi-llm-assistentov-v-kli?ysclid=mc2hvv983yr798737451> (дата обращения: 24.05.2025)
4. Beauchamp TL, Childress JF. Principles of biomedical ethics, fifth ed. New York: Oxford University Press. 2001; 454.
5. UN. Universal Declaration of Human Rights. New York, UN. [Internet]. Available from: <https://www.un.org/en/about-us/universal-declaration-of-human-rights>
6. Ayers JW. Internal Medicine Comparing Physician and Artificial Intelligence Chatbot Responses to Patient Questions Posted to a Public Social Media Forum. [Internet]. JAMA Internal Medicine Published online. 2023 April 28; 6. Available from: https://kstp.com/wp-content/uploads/2023/05/jamainternal_ayers_2023_oj_230030_1681999216.70842.pdf (accessed: 24.05.2025)
7. WHO. Ethics and governance of AI for health. [Internet]. 2023. Available from: <https://betterghanadigest.com/2023/05/17/who-calls-for-safe-and-ethical-ai-for-health/> (accessed: 24.05.2025)
8. Yanshen Sun. Exploring the Deceptive Power of LLM-Generated Fake News: A Study of Real-World Detection Challenges [Internet]. 2023; 16. Available from: <https://arxiv.org/html/2403.18249v1> (accessed: 24.05.2025)
9. Bloomberg. Economist Most Liveable Cities 2023 Ranking: Western Europe, Australia Top List. [Internet]. 2023. Available from: <https://www.bloomberg.com/news/articles/2023-06-22/economist-most-liveable-cities-2023-ranking-western-europe-australia-top-list>
10. Доминго Стивен. Health Service Journal (HSJ). [Internet]. 2023. Available from: <https://www.ccal.co.uk/post/three-takeaways-from-the-health-service-journal-hsj-digital-awards-2023> (accessed: 24.05.2025)

References

1. Kosarev Ye A. Uchebnik po mashinnomu obucheniyu. Available from: <https://education.yandex.ru/handbook/ml/article/yazykovye-modeli> (accessed: 24.05.2025). Russian.
2. Khokhlov AL, Kotlovskiy MYu, Pavlov AV, Potapov MP, Gabidullina LF, Tsybikova EB. Razvitiye neyrotekhnologiy: eticheskiye problemy i obshchestvennyye diskussii. Meditsinskaya etika. 2024; 1: 20–25. Russian.
3. Eticheskiye voprosy LLM-assistentov v klinike [Internet]. Meditsinskiye novosti i stat'i. 2025, may — [po sostoyaniyu na 18 iyunya 2025 goda]. Available from: <https://nexusacademy.ru/tpost/95kvl3v0k1-eticheskiye-voprosi-llm-assistentov-v-kli?ysclid=mc2hvv983yr798737451> (accessed: 24.05.2025) Russian.
4. Beauchamp TL, Childress JF. Principles of biomedical ethics, fifth ed. New York: Oxford University Press. 2001; 454.
5. UN. Universal Declaration of Human Rights. New York, UN. [Internet]. Available from: <https://www.un.org/en/about-us/universal-declaration-of-human-rights>
6. Ayers JW. Internal Medicine Comparing Physician and Artificial Intelligence Chatbot Responses to Patient Questions Posted to a Public Social Media Forum. [Internet]. JAMA Internal Medicine Published online. 2023 April 28; 6. Available from: https://kstp.com/wp-content/uploads/2023/05/jamainternal_ayers_2023_oj_230030_1681999216.70842.pdf (accessed: 24.05.2025)
7. WHO. Ethics and governance of AI for health. [Internet]. 2023. Available from: <https://betterghanadigest.com/2023/05/17/who-calls-for-safe-and-ethical-ai-for-health/> (accessed: 24.05.2025)
8. Yanshen Sun. Exploring the Deceptive Power of LLM-Generated Fake News: A Study of Real-World Detection Challenges [Internet]. 2023; 16. Available from: <https://arxiv.org/html/2403.18249v1> (accessed: 24.05.2025)
9. Bloomberg. Economist Most Liveable Cities 2023 Ranking: Western Europe, Australia Top List. [Internet]. 2023. Available from: <https://www.bloomberg.com/news/articles/2023-06-22/economist-most-liveable-cities-2023-ranking-western-europe-australia-top-list> (accessed: 24.05.2025)
10. Domingo Stephen. Health Service Journal (HSJ). [Internet]. 2023. Available from: <https://www.ccal.co.uk/post/three-takeaways-from-the-health-service-journal-hsj-digital-awards-2023> (accessed: 24.05.2025)