

ИСПОЛЬЗОВАНИЕ ОТКРЫТЫХ БОЛЬШИХ ЯЗЫКОВЫХ МОДЕЛЕЙ ДЛЯ БЫСТРОГО ПРОТОТИПИРОВАНИЯ СИСТЕМ ПЕРВИЧНОЙ ДИАГНОСТИКИ: СРАВНЕНИЕ С ТРАДИЦИОННЫМИ ML-ПОДХОДАМИ

Е.А. Макарова¹, А.В. Гусев²

¹ ООО «К-Скай», г. Петрозаводск, Россия;

² ФГБУ «Центральный научно-исследовательский институт организации и информатизации здравоохранения» Министерства здравоохранения Российской Федерации, г. Москва, Россия.

Поступила в редакцию: 25.05.2026

Принята к печати: 25.05.2026

Дата публикации: 30.06.2026

Для цитирования:

Макарова Е.А., Гусев А.В.
Использование открытых больших языковых моделей для быстрого прототипирования систем первичной диагностики: сравнение с традиционными ML-подходами. Организация и цифровизация здравоохранения, 2026. Т. 1. № 1. С. 22–27.

Контактная информация:

Гусев А.В.,
e-mail: agusev@webiomed.ai

Финансирование.

Исследование не имело спонсорской поддержки.

Конфликт интересов.

Авторы декларируют отсутствие явных и потенциальных конфликтов интересов в связи с публикацией данной статьи.

АННОТАЦИЯ

Актуальность. Внедрение систем поддержки принятия врачебных решений (СППВР) в практику первичной медико-санитарной помощи сопряжено со значительными временными и ресурсными затратами на этапах разработки и клинической валидации. Открытые большие языковые модели (LLM) с возможностью локального развертывания могут рассматриваться как инструмент быстрого прототипирования, позволяющий снизить издержки на начальных стадиях создания подобных систем, однако их использование имеет ряд ограничений.

Цель: провести сравнительный анализ временных и вычислительных затрат на разработку и эксплуатацию диагностических моделей на основе LLM и специализированных нейросетевых классификаторов (ML) для оценки экономической целесообразности применения LLM на этапе создания прототипа.

Материалы и методы. На выборке из 171 клинического случая с консенсусной разметкой, выполненной тремя врачами-экспертами, было проведено сравнение четырех квантованных LLM (DeepSeek-R17B и 14B, YandexGPT-5-Lite-8B, GigaChat-20B) и двух версий ML-модели Webiomed SC (обученных на 1 млн. и 18 млн. записей). Оценивались точность по метрике Top-3, время инференса, а также затраты человеко-часов и вычислительных ресурсов на подготовку и эксплуатацию моделей.

Результаты. Наилучшая из протестированных LLM (YandexGPT-8B) достигла точности 67% по метрике Top-3, что сопоставимо с ML-моделью, обученной на 1 млн. записей (67%), но уступает версии, обученной на 18 млн. записей (78%). При этом трудоемкость настройки LLM составила 8 человеко-часов против 340 человеко-часов для базовой ML-модели. Стоимость вычислительных ресурсов на этапе настройки LLM оказалась в 14 раз ниже затрат на обучение ML-модели на большом корпусе. Однако эксплуатационные расходы на LLM значительно превышают стоимость обслуживания ML-моделей.

Выводы. Открытые LLM выступают эффективным инструментом быстрого прототипирования при создании диагностических СППВР, позволяя в сжатые сроки и с минимальными начальными вложениями оценить жизнеспособность сервиса. Для промышленной эксплуатации в условиях ограниченных бюджетов региональных систем здравоохранения предпочтительным является использование специализированных ML-моделей, обеспечивающих низкую совокупную стоимость владения при высоких объемах обращений. Полученные данные могут быть использованы руководителями медицинских организаций при формировании стратегии цифровой трансформации диагностических процессов.

Ключевые слова: большие языковые модели, LLM, первичная диагностика, прототипирование, экономическая эффективность, машинное обучение, системы поддержки принятия врачебных решений, управление здравоохранением.

USING OPEN-SOURCE LARGE LANGUAGE MODELS FOR RAPID PROTOTYPING OF PRIMARY DIAGNOSIS SYSTEMS: A COMPARISON WITH TRADITIONAL ML APPROACHES

E.A. Makarova¹, A.V. Gusev²

¹K-Sky LLC, Petrozavodsk, Russia;

²Russian Research Institute of Health, Ministry of Health of the Russian Federation, Moscow, Russia.

Received: 25.05.2026

Accepted: 25.05.2026

Date of publication: 30.06.2026

For citation:

Makarova E.A., Gusev A.V.
Using open-source large language models for rapid prototyping of primary diagnosis systems: a comparison with traditional ML approaches. Healthcare Organization and Digitalization, 2026, Vol. 1. No. 1. P. 22–27.

For correspondence:

Gusev A.V.,
e-mail: agusev@webiomed.ai

Funding.

The study had no sponsorship.

Competing interests.

The authors declare the absence of any conflicts of interest regarding the publication of this paper.

ABSTRACT

Significance. The implementation of medical decision support systems (MDS) in primary healthcare practice is associated with significant time and resource costs at the stages of development and clinical validation. Open large language models (LLM) with the possibility of local deployment can be considered as a tool for rapid prototyping, which allows to reduce costs at the initial stages of creation of such systems, however, their use has a number of limitations.

Purpose: conduct a comparative analysis of the time and computational costs of developing and operating diagnostic models based on LLM and specialized neural network classifiers (ML) to assess the economic feasibility of using LLM at the prototype development stage.

Materials and methods. A comparison of four quantized LLMs (DeepSeek-R17B and 14B, YandexGPT-5-Lite-8B, GigaChat-20B) and two versions of the Webiomed SC ML model (trained on 1 million and 18 million records) was conducted on a sample of 171 clinical cases with consensus labeling performed by three expert doctors. The accuracy was evaluated using the Top-3 metric, as well as the inference time and the cost of human hours and computational resources for training and operating the models.

Results. The best LLM tested (YandexGPT-8B) achieved an accuracy of 67% using the Top-3 metric, which is comparable to the ML model trained on 1 million records (67%), but inferior to the version trained on 18 million records (78%). At the same time, the LLM configuration took 8 hours to complete, compared to 340 hours for the base ML model. The cost of computing resources during the LLM configuration phase was 14 times lower than the cost of training an ML model on a large corpus. However, the operational costs of LLM are significantly higher than the cost of maintenance.

Conclusions. Open LLMs are an effective tool for rapid prototyping when creating diagnostic support systems, allowing for a quick and cost-effective assessment of the service's viability. For industrial use in limited-budget regional healthcare systems, specialized ML models are preferred due to their low total cost of ownership and high volume of requests. The obtained data can be used by healthcare organizations' managers to develop a digital transformation strategy for diagnostic processes.

Keywords: large language models, LLM, primary diagnosis, prototyping, cost-effectiveness, machine learning, medical decision support systems, and healthcare management.

Введение

Цифровая трансформация здравоохранения предусматривает широкое внедрение интеллектуальных систем, способствующих повышению качества и оперативности диагностики при одновременном снижении нагрузки на медицинский персонал. Наиболее остро данная задача стоит в первичном звене, где дефицит времени и высокая вариабельность клинических проявлений создают риски диагностических ошибок.

Традиционный путь создания специализированных СППВР на основе методов машинного обучения требует значительных инвестиций: формирования и разметки массивов данных, обучения моделей, валидации, интеграции

в медицинские информационные системы. Для многих медицинских организаций, особенно регионального уровня, подобные затраты – как временные, так и финансовые – могут стать непреодолимым барьером на этапе принятия решения о внедрении.

Развитие больших языковых моделей и появление их открытых квантованных версий, пригодных для локального развертывания, открывает альтернативный путь – быстрое прототипирование диагностических сервисов с минимальными стартовыми вложениями. Такой подход позволяет в сжатые сроки проверить гипотезу о востребованности и точности сервиса, оценить пользовательский опыт и лишь затем принимать решение о целесообразности инвестиций

в разработку полноценного сервиса или медицинского изделия.

Цель настоящей работы: с позиций организатора здравоохранения оценить экономическую и временную эффективность использования открытых LLM для прототипирования систем первичной диагностики в сравнении с традиционным ML-подходом.

Обзор литературы

Важность поддержки первичной диагностики. Эффективность СППВР оценивается через влияние на принятие врачебных решений, сокращение ошибок и повышение качества медицинской помощи. Симптом-чекеры способствуют уменьшению времени постановки диагноза на ранних этапах и снижению риска пропуска важных симптомов, что доказано в том числе в российских исследованиях. Так, например, клинко-экономический анализ системы Webiomed. DHRA для оценки сердечно-сосудистого риска показал статистически значимое снижение 10-летней смертности (отношение шансов 0,33; 95% ДИ 0,27–0,40) [1]. Кроме того, Центр экспертизы и контроля качества медицинской помощи (ЦЭККМП) Минздрава России признал экономически приемлемым массовое внедрение ИИ-подсказок, прогнозируя сокращение риска смерти пациентов на 55,3% [2]. Внедрение таких инструментов в практику первичного звена может стать значимым резервом повышения эффективности использования ресурсов здравоохранения.

Технологические предпосылки: квантованные LLM. Квантование LLM является общепризнанным подходом для снижения вычислительных затрат [3]. В глобальном масштабе применение квантованных моделей рассматривается как ключевое направление обеспечения доступности ИИ-технологий для небольших медицинских организаций.

Экономические аспекты внедрения ИИ в здравоохранение. Концепция совокупной стоимости

владения (ТСО) для ИИ-решений включает в себя прямые затраты (на приобретение ПО, аппаратное обеспечение), косвенные (время персонала на обучение, изменение рабочих процессов) и скрытые (интеграция с существующими системами, нормативное соответствие). Принципиально важно, что «скрытые» затраты, связанные с очисткой и разметкой данных, обучением моделей, могут достигать 40% от общей ТСО [4]. Таким образом, необходимость инвестиций в этот этап должна быть экономически обоснована будущим применением, что ставит необходимость использования прототипирования до начала крупных разработок.

Материалы и методы

Подготовка данных и дизайн эксперимента. Для проведения сравнительного анализа был сформирован массив данных на основе реальных клинических случаев. Три врача-эксперта независимо оценивали корректность разметки диагнозов по МКБ-10 для каждого случая. В итоговую выборку вошел 171 случай, по которым было достигнуто полное согласие (консенсус) всех трех специалистов. Такой подход обеспечил высокое качество эталонной разметки, необходимое для объективного сравнения моделей.

Сравниваемые модели. В эксперименте использованы четыре открытые квантованные LLM, доступные для локального развертывания (таблица 1), и две версии специализированной ML-модели Webiomed SC2, обученной на различных объемах данных. Архитектура и подробности обучения первой версии модели описаны в работе [5].

ML-модели:

- 1. Webiomed SC2.0 – обучена на 1 млн. медицинских записей.
- 2. Webiomed SC2.1 – обучена на 18 млн. медицинских записей.

Выход моделей: уверенность (от 0 до 1) отнесения записи к определенному классу заболеваний, включающему несколько кодов МКБ.

Таблица 1

Характеристики использованных LLM

Название модели	Размер	Страна	Комментарии
deepseek-r1:7b	4,7 ГБ	Китай	Рассуждающая модель
deepseek-r1:14b [6]	9,0 ГБ	Китай	Рассуждающая модель
YandexGPT-5-Lite-8B-instruct-GGUF [7]	4,9 ГБ	Россия	Русскоязычная
GigaChat-20B-A3B-instruct-v1.5 [8]	14 ГБ	Россия	Русскоязычная

Таблица 2

Сравнительные результаты моделей

Модель	Топ-3 Accuracy	Время ответа (сек)
DeepSeek-7B	0,34	16,34
DeepSeek-14B	0,51	33,36
YandexGPT-8B	0,67	1,54
GigaChat-20B	0,56	1,05
Webiomed SC2.0 (1M)	0,67	0,2
Webiomed SC2.1 (18M)	0,78	0,2

В рамках методологии исследования использовался один промпт с целью унификации ответов языковой модели при анализе медицинских карт. Для повышения точности и единообразия выходного формата применялся few-shot подход [9]: модель знакомили с конкретными примерами корректного ответа (например, «J18.9 – Пневмония неуточненная – 85%»). От модели требовалось указать три наиболее вероятных диагноза, причём для каждого из них необходимо было привести код по МКБ-10, развёрнутое наименование заболевания и числовой процент надёжности. Таким образом, единый шаблон вывода позволял стандартизировать результаты вне зависимости от клинического случая.

При обработке ответом от LLM учитывалось не только совпадение по коду МКБ, но и текстовое описание диагноза, чтобы точнее привести к диагностическим группам и избежать ошибок, связанных с интерпретацией текстового ответа, который в подобной интерпретации существенно отличается от стандартного выхода ML-моделей.

Метрики и анализ затрат. Для оценки диагностической точности использовалась метрика Топ-3 Accuracy – доля случаев, в которых истинный диагноз присутствовал среди трех предсказанных моделью. Для экономического анализа фиксировались: время разработки (в человеко-часах), ресурсы на обучение/настройку и их стоимость (рассчитана по тарифам облачного провайдера для сопоставимых конфигураций), а также ресурсы на эксплуатацию и ежемесячная стоимость обслуживания при нагрузке, типичной для региональной поликлиники.

Результаты

Диагностическая точность и время выполнения. В таблице 2 представлены итоговые показатели точности и среднего времени инференса.

Лучшая из тестируемых LLM (YandexGPT-8B) продемонстрировала точность, идентичную ML-модели, обученной на 1 млн. записей, что свидетельствует о высоком качестве предобучения русскоязычных моделей. ML-модель, обученная на 18 млн. записей и адаптированная под конкретную клиническую задачу, ожидаемо превосходит все LLM, однако ее разработка потребовала на порядок больших ресурсов.

Стоит отметить, что подобная постановка задачи для LLM не раскрывает все сильные стороны этой технологии; при разработке LLM-системы, использующей внешние источники знаний, качество прогнозирования было бы вероятно выше, но в данном эксперименте оценивались возможности быстрого прототипирования.

Сравнительный анализ затрат на разработку и эксплуатацию. Ключевой интерес для организатора здравоохранения представляет сопоставление затрат на различных этапах жизненного цикла моделей. В таблице 3 приведены данные по трудозатратам и стоимости вычислительных ресурсов. В таблице указаны цены, актуальные на момент проведения эксперимента (июль 2025 года). Цены могут изменяться, но общий порядок разницы цен, вероятно, будет сохраняться.

Таблица 3

Затраты на разработку и эксплуатацию моделей

Реализация	Трудозатраты (чел.-час)	Ресурсы на настройку/ обучение	Стоимость настройки (руб.)	Ресурсы на эксплуатацию	Ежемесячная стоимость эксплуатации (руб.)
LLM (YandexGPT-8B)	8	GPU24GB, 72GB RAM, 16 CPU (24 ч)	2268	GPU24GB, 72GB RAM, 16 CPU	61440
Webiomed SC2.0	340	4GB RAM, 4 CPU	0*	4GB RAM, 4 CPU	3812
Webiomed SC2.1	168*	GPU24GB, 72GB RAM, 16 CPU (2 нед)	31752	4GB RAM, 4 CPU	3812

Примечания: * – для Webiomed SC2.1 указаны трудозатраты на дообучение и валидацию, без учета первоначальной разработки версии 2.0.

Интерпретация ключевых различий:

1. **Скорость прототипирования.** Настройка и валидация LLM заняла 8 человеко-часов, тогда как разработка базовой ML-модели потребовала 340 человеко-часов. Это означает, что гипотезу о практическом применении сервиса можно проверить за 1–2 рабочих дня вместо нескольких месяцев.
2. **Стоимость входа.** Настройка LLM обошлась на порядок меньше, чем обучение ML-модели на большом корпусе данных. Даже при рассмотрении только дообучения версии 2.1, разница составляет 14 раз.
3. **Стоимость эксплуатации.** Ежемесячное обслуживание LLM требует мощного GPU-сервера и оценивается примерно в 16 раз выше затрат на эксплуатацию легкой ML-модели. Это делает LLM неоптимальным решением для постоянной высоконагруженной эксплуатации в региональной поликлинике.

Обсуждение

LLM как инструмент быстрого прототипирования: управленческие выводы. Полученные результаты позволяют сформулировать стратегию развития диагностических СППВР на этапе проверки гипотез.

Этап 1. Быстрое прототипирование на LLM.

Использование открытой LLM позволяет в течение нескольких дней и с минимальными затратами создать работающий прототип системы. На этом этапе решаются следующие управленческие задачи:

- оценка востребованности сервиса врачами и/или пациентами;
- выявление типичных сценариев использования и узких мест в клинических процессах;
- сбор первичной обратной связи для уточнения требований к промышленной системе;
- демонстрация возможностей ИИ руководству и инвесторам для обоснования дальнейших инвестиций.

Этап 2. Переход на специализированный сервис. После подтверждения жизнеспособности продукта целесообразно инвестировать в разработку полноценного ИИ-сервиса (или дообучение специализированной ML-модели) с более подробной валидацией. Несмотря на значительные стартовые затраты (в основном на человеко-часы

аналитиков данных и медицинских экспертов), эксплуатационные расходы классических ML-моделей значительно ниже LLM, что обеспечивает быструю окупаемость при масштабировании.

Практические рекомендации для руководителей здравоохранения.

1. При оценке ИИ-решений учитывать совокупную стоимость владения на всем жизненном цикле. Низкая точность прототипа на «чистой» LLM может быть приемлема для проверки гипотез, но не для использования в реальной практике.
2. Использовать локальное развертывание LLM. Это позволяет избежать затрат на облачные API и, что более важно, обеспечивает соблюдение требований Федерального закона № 152-ФЗ «О персональных данных» при работе с медицинской информацией.

Ограничения исследования и направления дальнейшей работы.

- Выборка из 171 случая достаточна для пилотного сравнения, но для промышленной валидации требуется расширение массива данных.
- Не оценивалась стоимость интеграции с медицинскими информационными системами, которая может составлять значительную часть бюджета проекта.
- Перспективным направлением является оценка гибридных решений, сочетающих LLM для обработки сложных неструктурированных запросов и ML-модели для рутинной диагностики.
- Важно ответить, что в данном эксперименте используются далеко не все возможности LLM. Т.к. для систем характерны риски галлюцинаций и «потери» знаний при квантовании – в случае решения использования LLM в реальной практике целесообразно построение RAG систем и/или дообучение модели на целевых документах.

Выводы

Открытые квантованные LLM обеспечивают возможность быстрого прототипирования диагностических СППВР с минимальными временными и финансовыми затратами на этапе настройки. По точности Top-3 лучшая из протестированных LLM (67%) сопоставима с ML-моделью, обученной на 1 млн. записей, что достаточно для проверки концепции и сбора первичной обратной связи.

Эксплуатационные расходы на LLM существенно превышают затраты на обслуживание

специализированной ML-модели, что может сделать LLM экономически нецелесообразными для постоянной высоконагруженной эксплуатации. Для разработки подобных решений может быть использована двухэтапная стратегия: старт

с LLM-прототипа для быстрой проверки гипотез с последующим переходом на оптимизированную ML-модель при масштабировании сервиса или доработку LLM-решений для уменьшения риска галлюцинаций.

ЛИТЕРАТУРА

1. Морозов Д.Ю., Горкавенко Ф.В., Сeryapina Ю.В., Омеляновский В.В. Применение системы поддержки принятия врачебных решений Webiomed.DHRA при проведении диспансеризации в Российской Федерации: экономический анализ. Медицинские технологии. Оценка и выбор. 2025;47(3):61–73. <https://doi.org/10.17116/medtech2025470316>
2. Центр Минздрава признал экономически приемлемым массовое внедрение ИИ-подсказок для врачей // Медицинский вестник. 2024. URL: <https://medvestnik.ru/content/news/Centr-Minzdrava-priznal-ekonomicheskii-priemlemym-massovoe-vnedrenie-II-podskazok-dlya-vrachei.html> (Дата обращения: 04.06.2026).
3. Klang E., Apakama D., Abbott E.E. et al. A strategy for cost-effective large language model use at health system-scale. npj Digit. Med. 7, 320 (2024). <https://doi.org/10.1038/s41746-024-01315-1>
4. Cost of Implementing AI in Healthcare: A Decision-Maker's Guide [Электронный ресурс] // Signity Solutions. URL: <https://www.signitysolutions.com/blog/cost-of-implementing-ai-in-healthcare> (Дата обращения: 04.06.2026).
5. Ермак А.Д., Макарова Е.А., Кафтанов А.Н., Гаврилов Д.В., Новицкий Р.Э., Гусев А.В. Использование методов машинного обучения для диагностики заболеваний на основе неструктурированных медицинских текстов. Национальное здравоохранение. 2025;6(4):55–63. <https://doi.org/10.47093/2713-069X.2025.6.4.55-63>
6. Guo D., Yang D., Zhang H. et al. DeepSeek-R1 incentivizes reasoning in LLMs through reinforcement learning. Nature 645, 633–638 (2025). <https://doi.org/10.1038/s41586-025-09422-z>
7. YandexGPT-5-Lite-8B-pretrain [Электронный ресурс] // Hugging Face. – URL: <https://huggingface.co/yandex/YandexGPT-5-Lite-8B-pretrain> (Дата обращения: 04.06.2026).
8. GigaChat Family: Efficient Russian Language Modeling Through Mixture of Experts Architecture / V.Mamedov [et al.] // Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (ACL 2025). – 2025.
9. Brown T.B. Language Models Are Few-Shot Learners. (2020) arXiv: 2005.14165

REFERENCE

1. Morozov D.Yu., Gorkavenko F.V., Seryapina Yu.V., Omelyanovsky V.V. Application of the Webiomed.DHRA medical decision support system in conducting medical examinations in the Russian Federation: economic analysis. Medical technologies. Assessment and selection. 2025;47(3):61–73. <https://doi.org/10.17116/medtech2025470316> (In Russ.).
2. The Ministry of Health Center has recognized the mass implementation of AI prompts for doctors as economically feasible // Medical Bulletin. 2024. URL: <https://medvestnik.ru/content/news/Centr-Minzdrava-priznal-ekonomicheskii-priemlemym-massovoe-vnedrenie-II-podskazok-dlya-vrachei.html> (Accessed: 04.06.2026) (In Russ.).
3. Klang E., Apakama D., Abbott E.E. et al. A strategy for cost-effective large language model use at health system-scale. npj Digit. Med. 7, 320 (2024). <https://doi.org/10.1038/s41746-024-01315-1>
4. Cost of Implementing AI in Healthcare: A Decision-Maker's Guide // Signity Solutions. URL: <https://www.signitysolutions.com/blog/cost-of-implementing-ai-in-healthcare> (Accessed: 04.06.2026).
5. Ermak A.D., Makarova E.A., Kaftanov A.N., Gavrilov D.V., Novitsky R.E., Gusev A.V. Using Machine Learning Methods for Disease Diagnosis Based on Unstructured Medical Texts. National Healthcare. 2025;6(4):55–63. <https://doi.org/10.47093/2713-069X.2025.6.4.55-63> (In Russ.).
6. Guo D., Yang D., Zhang H. et al. DeepSeek-R1 incentivizes reasoning in LLMs through reinforcement learning. Nature 645, 633–638 (2025). <https://doi.org/10.1038/s41586-025-09422-z>
7. YandexGPT-5-Lite-8B-pretrain // Hugging Face. – URL: <https://huggingface.co/yandex/YandexGPT-5-Lite-8B-pretrain> (Accessed: 04.06.2026).
8. GigaChat Family: Efficient Russian Language Modeling Through Mixture of Experts Architecture / V.Mamedov [et al.] // Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (ACL 2025). – 2025.
9. Brown T.B. Language Models Are Few-Shot Learners. (2020) arXiv: 2005.14165

СВЕДЕНИЯ ОБ АВТОРАХ / ABOUT THE AUTHORS

Макарова Елена Андреевна – руководитель направления искусственного интеллекта, ООО «К-СКАЙ», г. Петрозаводск, Россия, e-mail: emakarova@webiomed.ru, ORCID ID 0000-0002-5410-5890

Elena A. Makarova – Head of the Artificial Intelligence Department, K-SKY LLC, Petrozavodsk, Russia, e-mail: emakarova@webiomed.ru, ORCID ID 0000 0002 5410 5890

Гусев Александр Владимирович – кандидат технических наук, эксперт по искусственному интеллекту, ФГБУ «Центральный научно-исследовательский институт организации и информатизации здравоохранения» Минздрава России, г. Москва, Россия; директор по развитию бизнеса, ООО «К-СКАЙ», г. Петрозаводск, Россия, e-mail: agusev@webiomed.ai, ORCID ID 0000-0002-7380-8460

Alexander V. Gusev – Ph.D. in Engineering, Expert in Artificial Intelligence, Russian Research Institute of Health, Moscow, Russia; Director for Business Development, K-SKAY LLC, Petrozavodsk, Russia, e-mail: agusev@webiomed.ai, ORCID ID 0000-0002-7380-8460