

# СРАВНИТЕЛЬНЫЙ АНАЛИЗ МЕТОДОВ ДООБУЧЕНИЯ LLM В УСЛОВИЯХ ОГРАНИЧЕННОСТИ ДЛЯ ДАННЫХ СПЕЦИАЛИЗИРОВАННЫХ ОБЛАСТЕЙ.

Нестерович Г. В., Фурсанов С. А.

Центр информатизации и инновационных разработок, факультет компьютерных систем и сетей,  
кафедра программного обеспечения информационных технологий,

Белорусский государственный университет информатики и радиоэлектроники

Минск, Республика Беларусь

E-mail: {g.nesterovich, s.fursanov}@bsuir.by

*В работе рассматриваются современные подходы к дообучению больших языковых моделей в условиях ограниченности данных в специализированных предметных областях. Проанализированы четыре метода: полное дообучение, параметр-эффективные техники, инструкционное дообучение и Retrieval-Augmented Generation. Для каждого подхода описаны принципы работы, выявлены преимущества и недостатки, а также показаны риски и ограничения, возникающие при дефиците размеченных данных.*

## ВВЕДЕНИЕ

Современные большие языковые модели (LLM) демонстрируют высокую универсальность и способность решать широкий спектр задач, однако при переносе на специализированные области возникает необходимость их дообучения. Проблема осложняется тем, что доступ к большим размеченным корпусам данных в медицине, праве, инженерии и других областях может быть ограничен по ряду причин: высокая стоимость разметки, наличие юридических и этических ограничений, конфиденциальность информации.

В таких условиях актуален вопрос выбора методов, позволяющих эффективно адаптировать модель при минимальных объёмах данных и ограниченных вычислительных ресурсах. Сравнительный анализ различных подходов к дообучению позволяет выявить их сильные и слабые стороны и предложить оптимальные стратегии для практического применения.

## I. Полное дообучение

Полное дообучение [1] подразумевает оптимизацию всех параметров предобученной модели на целевом наборе данных с задачами домена. Преимущество этого подхода состоит в максимальной выразительной мощности: модель может глубоко перестроить внутренние представления, что позволяет уловить сложные, нелинейные зависимости и тонкие семантические отличия, характерные для узкой предметной области. В случаях, когда доступны значительные по объёму и качеству экспертные аннотации, полное дообучение часто обеспечивает лучшие результаты по метрикам точности, полноты и соответствия форматам вывода. С другой стороны, при ограниченном объёме данных полный подход демонстрирует несколько серьёзных недостатков. Во-первых, высокой вероятностью выступает переобучение на небольшом и шумном наборе примеров, что ведёт к ухудшению обобщающей способности. Во-вторых, дорогостоящая вычислительная и про-

странственная нагрузка затрудняет экспериментальную итерацию и масштабирование в продуктивных системах: для каждой доменной версии требуется поддерживать либо отдельную копию модели, либо сложные механизмы переключения слоёв. В-третьих, полное дообучение сопряжено с риском катастрофического забывания общих навыков, что требует применения методов смешивания данных предобучения и доменных примеров, регуляризации, дистилляции или техник непрерывного обучения. Таким образом, полный подход оправдан в тех редких сценариях, когда есть обширные и качественные размеченные наборы и доступны значительные вычислительные ресурсы, но в условиях дефицита данных он требует дополнительных мер по защите от переобучения и высоких затрат.

## II. ПАРАМЕТР-ЭФФЕКТИВНОЕ ДООБУЧЕНИЕ

Параметр-эффективные методы (PEFT) [2] предлагают обучать небольшой набор дополнительных параметров или низкоранговых поправок, в то время как основная модель остаётся преимущественно замороженной. К типичным техникам относятся адаптеры, LoRA (low-rank adaptation), префикс-тюнинг и промт-тюнинг. Практическая ценность этих подходов заключается в значительном сокращении требуемой памяти и вычислительных затрат: для каждой предметной области достаточно хранить компактный пакет изменений вместо полной копии огромной модели, что облегчает развертывание и управление множеством доменных версий. При малых выборках PEFT демонстрируют повышенную устойчивость к переобучению, так как число обучаемых параметров невелико и риск подгонки под шум снижен. Эти методы хорошо сочетаются с техникой мультизадачного обучения и с передачей адаптеров между смежными областями, облегчают быстрое прототипирование и инкрементальные обновления. В то же время у PEFT есть и ограничения: из-за ограниченного числа новых

параметров теоретически может существовать потолок в том, какие представления можно эффективно скорректировать, и при задачах, требующих фундаментальной перестройки внутренних слоёв, маленькие адаптеры могут оказаться недостаточными. Дополнительно требуется аккуратный выбор мест внедрения адаптеров, ранга в LoRA и гиперпараметров обучения, на что может потребоваться эмпирическая валидация. Наконец, при агрегации множества адаптеров для сложных рабочих процессов возможно появление артефактов согласованности, которые надо отслеживать через метрики и человеческую оценку.

### III. ИНСТРУКЦИОННОЕ ДООБУЧЕНИЕ

Инструкционное дообучение [3] направлено на улучшение способности модели правильно интерпретировать и выполнять инструкции пользователя путём обучения на наборах пар «инструкция – желаемый ответ». Для специализированных доменов это особенно полезно, поскольку часто важнее соблюдение регламентов, форматов и последовательности рассуждений, чем просто воспроизведение фактов. При ограниченном количестве примеров аккуратная разработка инструкций, разнообразие формулировок и включение демонстраций форматов ответа могут существенно повысить полезность модели в рабочих сценариях: она начинает лучше соблюдать шаблоны, выдавать ответы требуемой структуры и избегать частых ошибок формата. Положительная сторона инструкционного подхода также в том, что он позволяет быстро корректировать поведение без глубокой перестройки весов модели, что экономит ресурсы и упрощает контроль качества. Минусы включают зависимость от качества самих инструкций: плохо сформулированные или синтетические инструкции способны ввести систематические смещения в поведение, а также не всегда обеспечивают добавление новых фактов или терминов в модельную базу знаний. Инструкционное дообучение часто лучше комбинировать с небольшой выборкой экспертных аннотаций и с элементами человеческой ревизии, чтобы сохранять корректность и избегать индуцированных ошибок.

### IV. RETRIEVAL-AUGMENTED GENERATION

Retrieval-Augmented Generation (RAG) [4] сочетает генеративную модель и модуль извлечения релевантных документов из внешней индексируемой базы знаний, подставляя найденный контекст в качестве дополнительного входа при генерации. Для узких доменов с ограниченной разметкой это крайне практичный путь: вместо того чтобы пытаться «вложить» все факты в веса модели, хранить и поддерживать внешнюю

базу, которую можно пополнять неразмеченными источниками, документами и справочниками. Преимущества RAG очевидны: оперативность обновления знаний без дорогостоящей переобучаемости, повышение фактологичности ответов при правильной настройке извлечения и возможность предоставления пользователю ссылок на источники и доказательств. Тем не менее эффективность RAG сильно зависит от качества индексирования, алгоритмов поиска и ранжирования, а также от полноты корпуса: если база неполна или содержит шум, генерация может ссылаться на нерелевантные или противоречивые фрагменты. Дополнительные вызовы включают необходимость верификации и агрегации противоречивых источников, обработку устаревших данных и обеспечение объяснимости ответов. В практических системах RAG часто комбинируют с локальными тонкими дообучениями или контролем качества на уровне ранжирования, применяют контрольно-верификационные пайплайны и меры по калибровке доверия.

### ЗАКЛЮЧЕНИЕ

Сравнение методов показывает, что при ограниченности данных оптимальные решения, как правило, гибридные и зависят от множества факторов: объёма и типа доступных текстов, наличия экспертной разметки, вычислительных и организационных ресурсов, требований к обновляемости знаний и к объяснимости. Полное дообучение остаётся мощным инструментом там, где доступны крупные качественные датасеты и ресурсы для проведения тщательной регуляризации и контроля, однако оно дорого и подвержено переобучению при малых выборках. Параметр-эффективные методы предлагают прагматичный компромисс, снижая затраты на хранение и развертывание и обеспечивая устойчивость при ограниченных примерах, при этом имея свои потолочные ограничения. Инструкционное дообучение повышает управляемость поведения модели и пригодность к форматированным задачам, но требует качественных инструкций и валидации. RAG даёт путь к обновляемости и фактической проверяемости, но требует аккуратной инженерии корпуса и поисковых компонентов.

1. Hands-On Large Language Models / J. Alammar, M. Grootendorst. // O'Reilly – 2024
2. Parameter-efficient fine-tuning of large-scale pre-trained language models / N. Ding et al. // Nature Machine Intelligence – 2023 – Vol. 5, P. 220–235
3. The Ultimate Guide to Fine-Tuning LLMs from Basics to Breakthroughs: An Exhaustive Review of Technologies, Research, Best Practices, Applied Research Challenges and Opportunities / V. Parthasarathy et al. // CeADAR Connect Group – 2024
4. A Simple Guide to Retrieval Augmented Generation / A. Kimothi // Manning – 2025