

Разработка и внедрение методов рассуждений в ИИ-ассистента ИВ РАН: эксперименты с «адаптивным мышлением» агента

On the development and implementation of reasoning methods for the IOS RAS AI assistant: conducting experiments with an agent's adaptive thinking capabilities

Гарин Артем Алексеевич

Artyom A. Garin

Кандидат исторических наук, руководитель Лаборатории цифровых исследований современного Востока Института востоковедения Российской академии наук

PhD (History), Head of the Digital Research Laboratory of Contemporary East at the Institute of Oriental Studies of Russian Academy of Sciences

a.garin@ivran.ru

ORCID: 0000-0003-4677-7221

В современной науке возросла необходимость более высокой интенсивности исследовательской деятельности, уровень академической конкуренции растет. Тенденция обусловлена широким внедрением методов искусственного интеллекта (ИИ) для автоматизации аналитики. ИИ-инструменты выполняют рутинные операционные задачи, помогают ученым наблюдать нестандартные корреляции и сложные паттерны в больших массивах данных, прежде ускользавшие из внимания при традиционном анализе. Ключевым моментом стало появление рассуждающих ИИ-помощников. Это позволяет формулировать более обоснованные гипотезы и проводить комплексный анализ, повышая значимость и качество ответов для специалистов. Прототип нового решения для построения интеллектуальных ассистентов создан на базе Лаборатории цифровых исследований современного Востока (ЛЦИСВ) ИВ РАН. Речь идет о более эффективной синхронизации цепочки рассуждений с возможностью обучения на лету. Новая версия ИИ-ассистента ИВ РАН может обучаться на собственном опыте (за счет накопления

Contemporary science is witnessing a growing demand for heightened research intensity, accompanied by an escalating level of academic competition. This trend is driven by the widespread integration of artificial intelligence (AI) methods for analytical automation. AI tools are capable of executing routine operational tasks and assist researchers in identifying non-standard correlations and complex patterns within large datasets, which previously eluded detection through traditional analysis. A pivotal development has been the emergence of reasoning AI assistants. This capability facilitates the formulation of more substantiated hypotheses and enables comprehensive analysis, thereby enhancing the relevance and quality of insights for specialists. A prototype of a novel solution for constructing such intelligent assistants has been developed at the Laboratory for Digital Research of the Contemporary East (LDRCE) within the Institute of Oriental Studies of the Russian Academy of Sciences (IOS RAS). The system focuses on achieving more efficient synchronization of the reasoning chain, incorporating the capacity for real-time learning. The new version of the IOS RAS AI assistant can learn from its own experience through the accumulation of knowledge by a

знаний контроллером), постепенно улучшая выдачу ответов на запросы пользователя.

dedicated controller, progressively improving the quality of its responses to user queries.

Ключевые слова: ИИ, рассуждения, граф знаний, нейросети, графовые нейронные сети, Азия

Keywords: AI, reasoning, knowledge graph, neural networks, GNN, Asia

Проблема «поверхностного» ИИ

Почему востоковеды стремятся научить ассистента не цитировать, а переосмысливать качественные данные?

Ключевой задачей в современной аналитике как в экономике, так и в международных отношениях остается ее качество и скорость получения. Быстрый и комплексный анализ, обоснованные прогнозы, ведущие к конкретным результатам, — основа для принятия решений и извлечения выгоды как в бизнесе, так и в науке. Объемы информации в истории, антропологии, этнографии, лингвистике и других гуманитарных сферах также колоссальны. Однако на сегодняшний день простого, базового «поверхностного» ответа на запрос к ИИ-ассистенту уже недостаточно. Более того, в отличие от точных показателей, качественные данные сложнее поддаются формализации и анализу. Наблюдается спрос на более обдуманную выдачу. Именно с этой проблемой — необходимостью осмысливать социо-гуманитарную информацию, выявлять в ней скрытые взаимосвязи и получать глубокие, продуманные выводы — и столкнулась рабочая группа на базе Лаборатории цифровых исследований современного Востока ИВ РАН.

Целью было создание новых версий ИИ-ассистента, способных анализировать исторические процессы, выстраивать логистические маршруты, разрабатывать стратегии делового сотрудничества. Иначе говоря, решать сложные задачи. Такой помощник должен оперировать даже специфическими деталями: учитывать тоннаж судов, просчитывать морские грузоперевозки за заданные промежутки времени и организовывать, к примеру, торговые пути из Юго-Восточной Азии в Россию.

Параллельно требовалось решить еще одну практическую задачу — более эффективно расходовать ресурсы, в данном случае токены, в которых тарифицируется работа больших языковых моделей. Для этого были оптимизированы режимы рассуждений.

Методология и обзор достижений в исследованиях рассуждающих ИИ-ассистентов

Создание размышляющих ИИ-ассистентов — давняя задача в области искусственного интеллекта, которая недавно наконец начала получать решения, отвечающие ее сложности. Сведения о новых открытиях появляются регулярно. Наблюдается сдвиг от теоретических изысканий к прикладным решениям. Выходят десятки статей, посвященных созданию размышляющих интеллектуальных помощников. В них рассматривается, как соединить многоагентную систему с нейросетью, чтобы агенты вместе учились и принимали простые решения [1]. Особенный интерес привлекают эксперименты с развитием «внимания» у моделей, их умением смотреть на важные составляющие входной информации [2].

Другие авторы предлагают простую «долгую память» для ассистентов, чтобы те помнили прошлые факты и диалоги без лишних затрат [4]. Среди прочих важных вопросов: как многоагентные системы чувствуют контекст и подстраиваются под обстановку [5] или показывают, как делать нейросети устойчивее к сбоям и атакам, чтобы ассистенты не ломались [6]. Иногда и вовсе случается, что длинные рассуждения не нужны — параллельно сгенерировать и проверить ответ бывает быстрее и точнее [14].

Уже описан переход от «быстрой интуиции» к «поэтапному мышлению» у больших языковых моделей [10]. Дело в том, что в «интуитивном» режиме базовые БЯМ работают как автодополнитель текста, выдают ответ одним проходом, почти не строя явных промежуточных рассуждений. Режим поэтапного мышления, напротив, предполагает разложение задачи на последовательность шагов: модель сначала планирует ход решения, затем генерирует цепочку рассуждений и затем формирует финальный ответ. Такой переход оказался востребован, потому что «быстрый» режим хорошо справляется с простыми, одношаговыми задачами, но дает много ошибок на сложных задачах.

Проведена систематизация подходов, которая показала, как несколько ИИ-ассистентов взаимодействуют и распределяют роли: в одних системах профили агентов заранее задаются разработчиком, в других — автоматически формируются самой моделью по описанию сценария или извлекаются из данных [8]. Вместе с тем изучением типов вывода и их «прокачкой» данными, обучением с подкреплением занимаются в том числе ученые из Университета Вестлейк, Чжэцзянского педагогического и Хайнаньского университетов [11]. Разбирались и способы рассуждать быстрее и дешевле — короче «думать», использовать меньшие модели и ускорять генерацию [6], или как ИИ-ассистенту раздавать подзадачи, спорить ради лучшего ответа, управлять памятью и внедрять всё это на практике [9]. Ученые из Чжэцзянского университета и компании Alibaba Group при этом учат ИИ-модели переключаться между режимами «быстро подумать» и «подумать основательно», получая лучшую точность при меньших тратах токенов [15].

Одним из самых значимых достижений стало появление агента Graph-R1, который строит лёгкий граф знаний и учится искать в нем и отвечать с помощью сквозного обучения с подкреплением [12]. Полную картину из чего вообще состоят ЛЛМ-агенты и где их применяют, дала команда специалистов из Китая, Сингапура и США [13].

Лидером по числу публикаций остается Китай. Названные работы оказали существенное влияние на разработку современных прикладных ИИ-решений. Рабочей группой на базе ЛЦИСВ ИВ РАН созданы три версии рассуждающих ИИ-ассистентов. Драйвером стало стремление добиться более глубокого и «осознанного» ответа от ИИ-модели, который бы устраивал ученых.

Для оптимизации процесса исследования и ускорения подбора эффективных комбинаций параметров кода, задействован в том числе генеративный искусственный интеллект. ИИ применялся как когнитивный усилитель в рамках контролируемого исследовательского процесса. На вход подавались детальные технические задания и критерии оптимизации, а его выводы подвергались критической оценке, корректировке ошибок и проверке. Основная цель применения ИИ заключалась в расширении исследовательских возможностей за счет его спо-

способности к быстрому перебору гипотез и многокритериальному анализу, служащих отправной точкой для дальнейшей экспертизы и экспериментальной проверки.

Эксперименты ЛЦИСВ ИВ РАН по автоматизации рассуждений ИИ-ассистента

В 2025 году в рамках разработки ИИ-ассистента ИВ РАН (работает на базе Yandex.Cloud [17]) сначала были созданы два базовых режима рассуждений, что соответствовало общемировым трендам в области искусственного интеллекта.

Рассуждающий режим **Reasoning (R.) AI Plus** (*предварительное название*) нацелен на полноту и качество ответа через промежуточные рассуждения. Иначе говоря, это средняя по «тяжести» версия рассуждающего ассистента. План ответа генерирует сама большая языковая модель (БЯМ). Для повышения надежности надстройка запрашивает у модели несколько вариантов такого плана (обычно состоящего из 3–5 пунктов), после чего выбирает самый полный и последовательно исполняет указанные шаги: сначала извлекает факты через генерацию, усиленную поиском (Retrieval-Augmented Generation, RAG), затем генерирует план из тезисов, потом проводит более глубокий анализ идей или вычислительные проверки и наконец сводит результаты в единый ответ. Чтобы «придумать» идеи, R. AI Plus строит древо рассуждений (Tree of Thought, ToT), в каждом узле которого может попробовать, например, разные варианты, а потом выбрать лучший. Если нужно что-то быстро посчитать или проверить, она использует программный метод мышления (Program-of-Thought, PoT) – запускает небольшой фрагмент кода на Python и получает точный результат. Этот подход обеспечивает более стройную аргументацию и опору на фактические данные, однако генерация самого плана, особенно при многократной попытке подбора наилучшего варианта, требует дополнительных токенов на сам процесс планирования и на форматирование ответа, что увеличивает суммарный расход токенов и задержку по сравнению с одношаговым ответом.

Режим **Reasoning AI Sprint** (*предварительное название*) — это облегченная рассуждающая версия, основная задача которой — дать максимально быстрый ответ с минимальными затратами. При получении запроса система анализирует ключевые слова, подбирает наиболее релевантный инструмент (например, модуль извлечения данных для поиска фактов или «мозговой штурм»), выполняет его и выдает результат. Весь процесс сводится к двум обращениям: к инструменту (если такой удалось подобрать) и к самой языковой модели. Это позволило достичь минимальной задержки и экономии токенов (в сравнении с «тяжелым» рассуждающим ассистентом), но сказывается на полноте ответа. R. AI Sprint подходит для простых вопросов и интеграции в сопутствующие приложения, где важна скорость и легкость выдачи.

Появление нового метода Graph-R1 и разработка режима Reasoning GraphMind в ИИ-ассистенте ИВ РАН

В течение 2025 г. планировалось развивать и анализировать два режима рассуждений ИИ-ассистента ИВ РАН, после чего подготовить итоговую статью. Однако динамика исследований в сфере искусственного интеллекта буквально вынуждает менять планы из-за регулярных технологических скачков и появления инновационных подходов. 30 июля 2025 г. на сайте препринтов arXiv.org вышла

статья «*Graph-R1: Towards Agentic GraphRAG Framework via End-to-end Reinforcement Learning*» [12], авторы которой предложили новый подход к генерации, усиленной поиском (RAG). Вместо одноразового извлечения фрагментов текста ученые из Китая предложили формировать гиперграф знаний и рассматривать процесс поиска и генерации как многошаговое взаимодействие агента с окружением. Иначе говоря, гиперграф — структура, где одно «гиперребро» может соединять сразу несколько сущностей. Это позволяет более точно передавать смысл сложных предложений. Модель работает как агент, который не просто один раз извлекает нужные данные, а может по шагам "думать, переформулировать запрос, извлекать новый контекст, снова думать — и только потом отвечать. Вся эта логика обучается от начала и до конца с помощью специально разработанного метода обучения с подкреплением под названием «групповая относительная оптимизация политики» (Group Relative Policy Optimization, GRPO), где для одного запроса генерируется несколько ответов, они сравниваются между собой, и агент получает относительное вознаграждение за более корректные, структурированные и логически обоснованные решения. Это позволяет не только снизить число нерелевантных запросов, но и повысить точность обоснований и качество сгенерированного текста. Эксперименты на стандартных RAG-датасетах показали, что Graph-R1 превосходит классические GraphRAG-методы и другие ускоренные RAG-подходы по точности, эффективности и качеству итоговых ответов.

Это достижение подтолкнуло к череде экспериментов, проведенных на базе Лаборатории цифровых исследований современного Востока ИВ РАН. Рабочая группа поставила перед собой ряд задач. Во-первых, облегчить этот механизм, чтобы использовать не только на мощной инфраструктуре, а сделать доступным для среднего «железа» и ИИ-ассистентов. Во-вторых, развитие получило внедрение графовых нейронных сетей (*Graph Neural Networks, GNN*), чтобы сделать рассуждение ИИ-ассистента более «динамичным». GNN пропускают информацию по графу, обучая представления каждого узла с учетом его окружения (состояние узлов обновляется в несколько проходов).

Такой режим рассуждающего ИИ-ассистента получил название R. AI GraphMind. Это «умный» сценарий размышлений, который на каждом шаге решает, что делать дальше, опираясь сразу на большую языковую модель и маленькую вспомогательную нейросеть. Последняя смотрит на краткое представление текущей ситуации и предлагает действие; итоговый выбор — это смесь ее предложения и совета большой модели, с небольшой долей случайного «исследования». Если становится понятно, что дальше думать не нужно, система досрочно завершает рассуждения и выдаёт ответ.

Процесс осуществляется следующим образом: сначала большая модель набрасывает короткий план. Затем система по шагам либо сразу формирует ответ, либо делает быстрый вычислительный ход (генерирует фрагмент Python и исполняет его), либо устраивает «мозговой штурм» — просит у модели несколько независимых идей. При необходимости подмешиваются факты из веб-поиска. Если новые шаги не дают пользы, Reasoning AI GraphMind останавливается и собирает финальный ответ. Чтобы не раздувать контекст повторами, в память не добавляются почти одинаковые куски, что экономит место. Маленькая сеть дообучается, «подражая» БЯМ — она учится выбирать так же, как выбирала боль-

шая модель на прошлых шагах. На простых запросах R. AI GraphMind чаще ограничивается короткими ходами и отвечает несколько быстрее. На сложных — добавляет «мозговой штурм» и/или вычисления и обычно дает более продуманный ответ, на который требуется больше времени. Количество затраченных токенов и временная задержка зависят от того, сколько вспомогательных шагов действительно потребовалось.

Вероятно, подобный механизм позволит снизить и риски галлюцинаций. Модель может меньше сбиваться, когда не нужно повторно генерировать похожие тексты. Чем всегда выполнять и «мозговой штурм», и расчеты, R. AI GraphMind может пропустить лишний шаг. Предыдущая версия, получив план от БЯМ, выполняла все узлы плана, даже если какие-то шаги могли оказаться избыточными. Новая версия пропускает ненужные действия: это снижает среднее число шагов без потери качества решения.

Однако впоследствии возникли опасения, что подобный подход негативно скажется на креативности ИИ-ассистента, который будет выдавать практически одинаковые ответы, вместо более глубоких выводов. Предприняв ряд шагов, удалось предотвратить жесткого следования выученной структуре, которое бы сделало ответы более однотипными. Например, при каждом срабатывании узла принятия решения R. AI GraphMind с вероятностью 10% и более вместо использования выученного правильного ответа система уходит на «исследование» и делает запрос к ИИ-модели. Кроме того, можно накапливать примеры в формате «контекст-действие», прежде чем переводить узел на «обученный» поиск и принимать решения без обращения к БЯМ. Иначе говоря, система запоминает, в какой ситуации она прежде была и что в этой ситуации решила делать. GNN-контроллер смотрит на накопленное множество пар «контекст – выбранное действие» и выбирает следующий шаг. Там, где GNN стабильно угадала много раз, узел (или кластер похожих состояний) можно считать «обученным». Можно изменять и температуру рассуждений.

Первые сравнения рассуждающих режимов ИИ-ассистента ИВ РАН (R. AI GraphMind, R. AI Plus и R. AI Sprint)

Первые тесты режимов рассуждений ИИ-ассистента проводились на этапе разработки и носили ограниченный характер. Далее будут представлены предварительные результаты, которые показали работоспособность подходов, однако для полноценной оценки их эффективности все равно потребуется больше информации.

Для работы ИИ-ассистента используется модель-эмбеддер paraphrase-multilingual-MiniLM-L12-v2 [16], которая превращает текст в числовые векторы. В режиме R. AI GraphMind, как подчеркивалось, поверх основной модели работает графовый модуль. Для сравнения разных режимов Reasoning AI использовались языковые модели, доступные в сервисе Yandex.Cloud (YandexGPT и YandexGPT Pro). В ходе проверки функционала ИИ-ассистенту ИВ РАН подавались запросы о странах Востока.

Первые тесты показали, что режим R. AI GraphMind сохраняет или повышает качество ответов относительно R. AI Plus и R. AI Sprint при сложных задачах. На базе Лаборатории цифровых исследований современного Востока ИВ РАН сравнили три рассуждающих режима ассистента на типовых заданиях (составление

аналитических текстов, прогнозов и тем для научных исследований). Ниже приведена сводная таблица с предварительными оценками качества ответа (оценивалось экспертом ЛЦИСВ ИВ РАН), временем ответа и количества затраченных токенов (предварительная оценка).

Рассуждающий режим	Ожидаемое качество (из трех)	Время ответа	Затраченные токены
Reasoning AI Plus	Хорошее	Среднее (14–18 сек.)	≈2300
Reasoning AI Sprint	Среднее	Самое быстрое из трех (4–8 сек.)	≈1900
Reasoning AI GraphMind	Высокое	Медленное (21–25 сек.)	≈2500

Таблица 1. Сравнение эффективности рассуждающих режимов ИИ-ассистента ИВ РАН

R. AI GraphMind показал себя как инструмент для выработки стратегических предложений. Он находит баланс между Plus и Sprint, но требует больших временных затрат.

Наблюдалась разница и в количестве затраченных токенов между тремя режимами ИИ-ассистента ИВ РАН. R. AI GraphMind использовал больше токенов, чем Plus и Sprint. Объем выходного текста как параметр был схожим, однако внутренние процессы их генерации и конечный результат различаются.

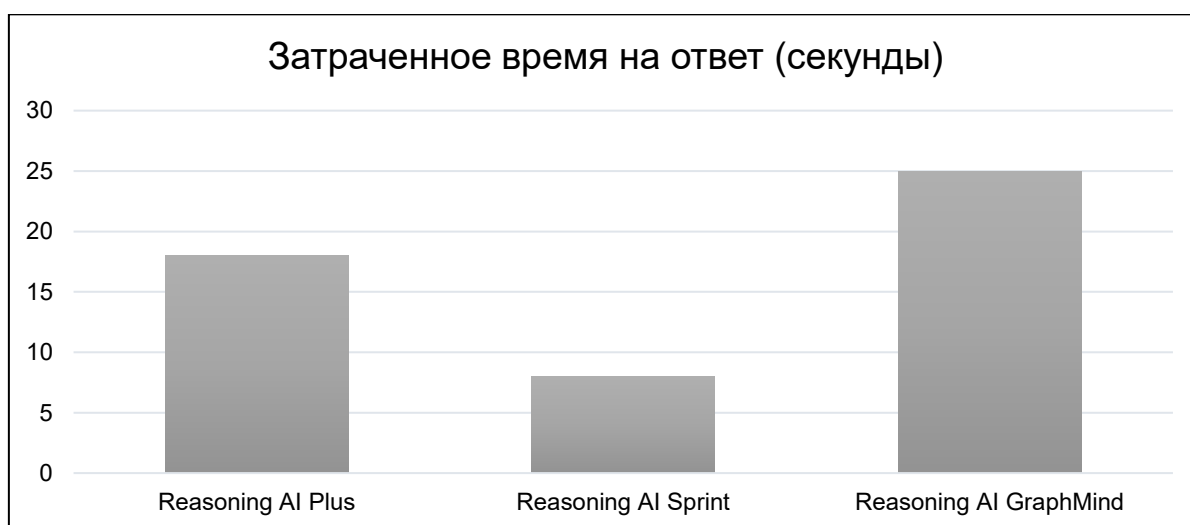


Рис. 1. Результаты измерений времени от отправки запроса до готового ответа. Составлено автором на основе механизма телеметрии, внедренного в рассуждающие режимы ИИ-ассистента ИВ РАН.

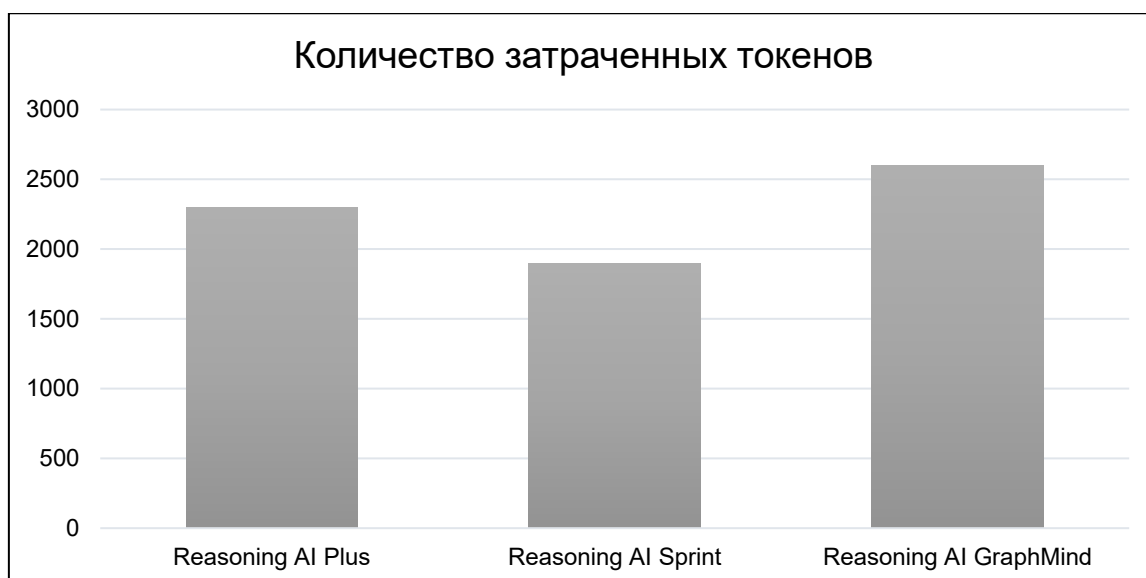


Рис. 2. Результаты измерений количества затраченных токенов от отправки запроса до готового ответа. Составлено автором на основе механизма телеметрии, внедренного в рассуждающие режимы ИИ-ассистента ИВ РАН.

Рассмотрим конкретный пример. При тестировании прототипа особое внимание уделялось изучению логистической связности в Юго-Восточной Азии. Все три режима проверялись запросом «Порты Юго-Восточной Азии», чтобы проверить, как они выстроят свой ответ, если не будет задано конкретных инструкций. Ответ R. AI GraphMind оказался наиболее релевантным из трех. Перечисленные им пять портов действительно географически относятся к Юго-Восточной Азии. Сингапур, Джакарта (Танджунг Приок), Лаем Чабанг в Таиланде и Хошимин во Вьетнаме — это ключевые морские логистические узлы региона. Однако порт Келанг в Малайзии в ответе представлен в первую очередь как «пассажирский круизный центр», хотя он выступает одним из самых загруженных грузовых контейнерных портов в мире. R. AI GraphMind также оказался единственным режимом, который сопровождал список портов небольшой справкой: *«Эти порты оснащены современной инфраструктурой и способны обрабатывать большие объёмы грузов, включая контейнеры, нефть и другие товары. Географическое положение портов Юго-Восточной Азии способствует их значимости как транспортных узлов, соединяющих Европу, Азию и Африку морскими путями. Развитие портов в регионе способствует экономическому росту и созданию рабочих мест, что делает их важными центрами для местных и международных инвесторов»*.

Ответ R. AI Plus оказался самым кратким, но при этом все три названных порта — Сингапур, Танджунг Приок и порт Келанг — абсолютно верны и выступают ядром портовой системы региона. При этом излишняя краткость одновременно стала его минусом. R. AI Sprint выдал ответ быстрее всех. Он предоставил список сразу из 11 портов, но допустил географические недочеты. R. AI Sprint включил в перечень портов Шанхай, Гонконг, Пусан и Токио, которые действительно являются одними из крупнейших в мире, но они расположены в Восточной Азии, а не в Юго-Восточной.

R. AI GraphMind также показал способность к более глубокой автономной работе с внутренней базой знаний, извлекая и структурируя информацию о портах

даже без доступа к веб-источникам, в то время как R. AI Sprint проявил большую predisposedность к работе с данными из сети.

Перспективы графового режима Reasoning AI GraphMind

Появление Graph R-1 стало важным этапом в эволюции искусственного интеллекта. Это новый комплексный взгляд, как модели могут взаимодействовать с внешними источниками знаний через RAG-архитектуру и оптимизировать свои действия. Работа китайских специалистов стала вдохновением для ЛЦИСВ ИВ РАН: она показала, что единый агент способен не просто последовательно извлекать информацию и генерировать текст, а выстраивать динамический гиперграф, по которому управляет своим поиском и формирует обоснованные ответы.

На базе лаборатории выдвинут ряд гипотез об оптимизации размышлений, а также создан прототип режима рассуждений ИИ-ассистента ИВ РАН под названием R. AI GraphMind. По сравнению с R. AI Plus, новый режим достигает более высокого качества ответов. А использование GNN стало шагом к «адаптивному мышлению» ИИ-ассистента ИВ РАН.

Однако это не значит, что каждый из режимов рассуждений будут использоваться лишь в одном чат-боте. Например, R. AI GraphMind можно потенциально использовать, чтобы контролировать сложные процессы или сформировать более глубокий отчет по документу, решая сложные задачи. R. AI Plus может подойти для более креативных запросов. В веб-приложениях перспективным выглядит режим Reasoning AI Sprint для быстрого отклика, в том числе при работе кнопок в различных сервисах.

Комбинирование способностей БЯМ с графовыми методами — крайне перспективное направление исследований, создания более умных и экономичных ИИ-систем, способных рассуждать и, учиться на лету.

Библиография | References

1. Asadi, R., Mustapha, N., & Sulaiman, N. A framework for intelligent multi agent system based neural network classification model. 2009. *arXiv preprint arXiv:0910.2029*
2. Bahdanau, D., Cho, K., & Bengio, Y. *Neural Machine Translation by Jointly Learning to Align and Translate*. 2015. <https://arxiv.org/abs/1409.0473>
3. Casper, S., Bailey, L., Hunter, R., Ezell, C., Cabalé, E., Gerovitch, M., Slocum, S., Wei, K., Jurković, N., Khan, A., Christoffersen, P., Özişik, A. P., Trivedi, R., Hadfield-Menell, D., & Kolt, N. The AI Agent Index. 2025. <https://arxiv.org/html/2502.01635v1>
4. Chhikara, P., Khant, D., Aryan, S., Singh, T., & Yadav, D. Mem0: Building Production-Ready AI Agents with Scalable Long-Term Memory. 2025. <https://arxiv.org/abs/2504.19413>
5. Du, H., Thudumu, S., Nguyen, H., Vasa, R., & Mouzakis, K. A comprehensive survey on context-aware multi-agent systems: Techniques, applications, challenges and future directions. 2025. <https://arxiv.org/html/2402.01968v2>
6. Duddu, V., Pillai, N. R., Rao, D. V., & Balas, V. E. (2020). Fault tolerance of neural networks in adversarial settings. *Journal of Intelligent & Fuzzy Systems*, 38(5), 5897–5907. <https://doi.org/10.3233/JIFS-179677>
7. Feng, S., Fang, G., Ma, X., & Wang, X. Efficient Reasoning Models: A Survey. 2025. <https://arxiv.org/html/2504.10903v1>

8. Guo, T., Chen, X., Wang, Y., Chang, R., Pei, S., Chawla, N. V., Wiest, O., & Zhang, X. Large Language Model based Multi-Agents: A Survey of Progress and Challenges. 2024. <https://arxiv.org/html/2402.01680v2>
9. Han, S., Zhang, Q., Yao, Y., Jin, W., & Xu, Z. LLM Multi-Agent Systems: Challenges and Open Problems. 2025. <https://arxiv.org/abs/2402.03578>
10. Li, Z.-Z., Zhang, D., Zhang, M.-L., Zhang, J., Liu, Z., Yao, Y., Xu, H., Zheng, J., Wang, P.-J., Chen, X., Zhang, Y., Yin, F., Dong, J., Li, Z., Bi, B.-L., Mei, L.-R., Fang, J.-F., Liang, X., Guo, Z., Song, L., & Liu, C.-L. From System 1 to System 2: A Survey of Reasoning Large Language Models. 2025. <https://arxiv.org/abs/2502.17419>
11. Liu, H., Fu, Z., Ding, M., Ning, R., Zhang, C., Liu, X., & Zhang, Y. Logical Reasoning in Large Language Models: A Survey. 2025. <https://arxiv.org/abs/2502.09100>
12. Luo, H., E, H., Chen, G., Lin, Q., Guo, Y., Xu, F., Kuang, Z., Song, M., Wu, X., Zhu, Y., & Tuan, L. A. Graph-R1: Towards Agentic GraphRAG Framework via End-to-end Reinforcement Learning. 2025. <https://arxiv.org/abs/2507.21892>
13. Luo, J., Zhang, W., Yuan, Y., Zhao, Y., Yang, J., Gu, Y., Wu, B., Chen, B., Qiao, Z., Long, Q., Tu, R., Luo, X., Ju, W., Xiao, Z., Wang, Y., Xiao, M., Liu, C., Yuan, J., Zhang, S., Jin, Y., Zhang, F., Wu, X., Zhao, H., Tao, D., Yu, P. S., & Zhang, M. Large Language Model Agent: A Survey on Methodology, Applications and Challenges. 2025. <https://arxiv.org/abs/2503.21460>
14. Ma, W., He, J., Snell, C., Griggs, T., Min, S., & Zaharia, M. Reasoning Models Can Be Effective Without Thinking. 2025. <https://arxiv.org/abs/2504.09858>
15. Xiao, W., Gan, L., Dai, W., He, W., Huang, Z., Li, H., Shu, F., Yu, Z., Zhang, P., Jiang, H., & Wu, F. Fast-Slow Thinking for Large Vision-Language Model Reasoning. 2025. <https://arxiv.org/html/2504.18458v1>

Программное обеспечение и документация | Software and Manuals

16. Paraphrase-multilingual-MiniLM-L12-v2 : [программное обеспечение]. – 2020. – URL: <https://huggingface.co/sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2>.
17. Yandex Cloud. Foundation Models : [электронный документ]. – 2023. – URL: <https://yandex.cloud/en/docs/foundation-models/>