

КОМПЬЮТЕРНОЕ МОДЕЛИРОВАНИЕ COMPUTER SIMULATION

УДК 004.85

DOI 10.52575/2687-0932-2026-53-2-388-399

EDN OJWXHS

Сравнительный анализ векторных, графовых и гибридных моделей представления знаний в системах информационного поиска

Настасенко С.А., Савотченко С.Е.

Российский государственный геологоразведочный университет им. Серго Орджоникидзе,
Россия, Москва, 117997, ул. Миклухо-Маклая, д. 23
snastasenko99@gmail.com, savotchenkose@mgri.ru

Аннотация. Представлены результаты анализа современных подходов к информационному поиску, основанных на векторных и графовых моделях представления знаний. Выявлены преимущества и ограничения каждого подхода, включая вопросы интерпретируемости, масштабируемости и полноты знаний. Особое внимание уделяется гибридным методам, сочетающим семантическую глубину векторных представлений с точностью структурированных графовых знаний. Показано, что гибридные подходы демонстрируют высокий потенциал для повышения релевантности поиска, однако они требуют дальнейшего развития универсальной методологии для интеграции гетерогенных представлений. На основе проведенного анализа сформулированы практические рекомендации, направленные на развитие и совершенствование гибридных моделей.

Ключевые слова: информационный поиск, векторные модели, графы знаний, графовые модели, семантические эмбединги, гибридные модели, нейросетевые архитектуры

Для цитирования: Настасенко С.А., Савотченко С.Е. 2026. Сравнительный анализ векторных, графовых и гибридных моделей представления знаний в системах информационного поиска. *Экономика. Информатика*, 53(2): 388–399. DOI 10.52575/2687-0932-2026-53-2-388-399. EDN OJWXHS

Comparative Analysis of Vector, Graph and Hybrid Models of Knowledge Representation in Information Retrieval Systems

Sergey A. Nastasenko, Sergey E. Savotchenko

Sergo Ordzhonikidze Russian State University for Geological Prospecting,
23 Miklukho-Maklay St., Moscow 117997, Russia
snastasenko99@gmail.com, savotchenkose@mgri.ru

Abstract. The paper presents the results of an analysis of modern approaches to information retrieval based on vector and graph models of knowledge representation. The advantages and limitations of each approach are identified, including issues of interpretability, scalability, and knowledge completeness. Particular attention is given to hybrid methods that combine the semantic depth of vector representations with the precision of structured graph knowledge. It is shown that hybrid approaches demonstrate high potential for improving search relevance; however, they require further development of a universal methodology for integrating heterogeneous representations. The authors provide practical recommendations for the development and improvement of hybrid

models based on the analysis. The study also emphasizes that the effectiveness of such systems depends on data preprocessing, consistent links between entities and the choice of ranking algorithms. These conclusions can support the design of search systems that require deep contextual understanding.

Keywords: information retrieval, vector models, knowledge graphs, graph models, semantic embeddings, hybrid models, neural network architectures

For citation: Nastasenko S.A., Savotchenko S.E. 2026. Comparative Analysis of Vector, Graph and Hybrid Models of Knowledge Representation in Information Retrieval Systems. *Economics. Information technologies*, 53(2): 388–399 (in Russian). DOI 10.52575/2687-0932-2026-53-2-388-399. EDN OJWXHS

Введение

Современный информационный поиск сталкивается с постоянно растущими объемами данных и всё более сложными структурами информации. Пользователи формулируют запросы на естественном языке, ожидая получить релевантные ответы из массивов документов, социальных сетей, мультимедийных данных и других источников [Kruhlov, 2024]. Главной проблемой является то, как представить знания, содержащиеся в этих неструктурированных данных, в форме, пригодной для эффективного поиска. Представление знаний в информатике занимается формализацией и структурированием сведений (например, в виде онтологий, баз знаний или графов знаний) с целью их дальнейшего использования алгоритмами [Kyurkchiev, 2026]. Однако традиционные методы поиска долгое время опирались преимущественно на статистические модели текста, не учитывая явные семантические связи между объектами, которые представлены в базах знаний.

Прямое применение графовых баз знаний для ответа на произвольные пользовательские запросы сопряжено с трудностями: требуется предварительно преобразовать текст запроса в формальные понятия и отношения, а также обеспечить покрытие знаниями всех возможных вопросов. Далеко не каждый утилитарный запрос пользователя легко переводится на язык онтологии или SPARQL-запроса, особенно если формулировка неполная или двусмысленная. Кроме того, базы знаний по своей природе ограничены тем набором фактов, который в них заложен, и могут не содержать новейшей или узкоспециализированной информации, которая имеется в текстовых данных.

Объём и сложность данных также усугубляют проблему интеграции знаний в поиск [Chang, 2024]. Создание всеобъемлющих графов знаний, охватывающих абсолютно всю информацию из постоянно обновляющихся источников, практически невыполнимо из-за чрезвычайной трудоёмкости и динамичности данных. Даже крупнейшие онтологии и базы знаний (например, Wikidata или Knowledge Graph от Google) не успевают в реальном времени пополняться всеми новыми фактами. Поэтому на практике приходится совмещать подходы: векторные модели эффективно обрабатывают неструктурированные тексты и учатся на больших корпусах данных [Chatterjee, 2025], в то время как графовые модели вносят структурные, семантически значимые связи между сущностями [Chao Li., 2025].

Возникает потребность в гибридных методах, которые могут преодолеть упомянутый семантический разрыв, соединяя статистическую мощь векторных представлений с экспрессивностью графовых знаний [Parageorgiou G, 2025]. Проблематика информационного поиска сегодня во многом сводится к вопросу: как интегрировать явные знания (например, факты и отношения из графов) в процесс поиска по неструктурированным данным, сохраняя при этом высокую масштабируемость и точность поиска [Parageorgiou G, 2025].

Векторные методы получили широкое распространение благодаря своей простоте и высокой производительности, особенно с развитием нейросетевых моделей и трансформеров [Абрамович, 2025]. Параллельно развивались графовые подходы – от классических алгоритмов обхода графов до современных графовых нейросетей (GNN) и графов знаний (KG). [Zhu, 2024].

В этом контексте особую актуальность приобретают гибридные алгоритмы, сочетающие сильные стороны графовых и векторных представлений [Настасенко, 2025]. Такие методы открывают возможности более точного моделирования смысла, устойчивости к шуму,

повышения интерпретируемости результатов, а также адаптации под конкретные предметные области. Разработка новых гибридных алгоритмов поиска представляет собой актуальное направление в области интеллектуального анализа данных, искусственного интеллекта и прикладной информатики [Sarmah, 2024].

Объект и методы исследования

Объектом данного исследования являются процессы поиска информации в контексте неструктурированных текстовых массивов и структурированных баз знаний. Предметом исследования являются векторные, графовые и гибридные модели (методы и алгоритмы представления информации). Важное место занимают способы интеграции векторных и графовых подходов для повышения эффективности, точности и интерпретируемости поиска [Настасенко, 2025]. В работе рассматриваются закономерности формирования семантических связей между сущностями, специфика их формализации в графах знаний и потенциал для сочетания статистических и логико-символических методов обработки информации

Методологической основой исследования является системный анализ, позволяющий рассматривать поиск информации как сложную задачу, охватывающую лингвистические, вычислительные и когнитивные аспекты. Теоретическая основа работы базируется на фундаментальных положениях теории информации, математической лингвистики и теории графов, а также на современных концепциях представления знаний в интеллектуальных системах. В исследовании используются методы сравнительного анализа для выявления преимуществ и ограничений различных подходов к поиску информации. Модели сравнивались по их пользе для поисковых систем. Важно было понять, как каждая модель хранит знания. Также учитывалось, можно ли объяснить результат поиска. Отдельно рассматривалась работа с большими объемами данных. Еще учитывалось, умеет ли модель использовать логические связи. Также оценивалась работа с обычным неструктурированным текстом. Такой анализ показывает, где каждая модель работает хорошо, а где у нее есть ограничения.

В работе также показано, как менялись методы поиска. Сначала использовались статистические модели. Они опирались на частоту слов и совпадение терминов. Потом стали развиваться нейросетевые методы. Они лучше учитывают смысловую близость текстов. Сейчас всё чаще используются гибридные решения. В них векторный поиск дополняется графами знаний. Поэтому развитие поиска идет не через полный отказ от старых методов. Скорее, разные подходы постепенно объединяются в одну систему.

Результаты и обсуждение

Модели поиска: векторные, графовые и гибридные

Одним из первых и наиболее распространённых подходов к поиску стала векторная модель, в которой каждый документ описывается набором признаков (например, словами или терминами), которым соответствуют координаты вектора. Преимуществами векторных моделей являются высокая обобщающая способность и умение работать с неструктурированным текстом без ручной разметки знаний. Они автоматически улавливают статистические закономерности языка, что особенно хорошо масштабируется на больших данных [Sadykova, 2025]. Однако есть и недостатки: векторные представления плохо интерпретируемы (сложно объяснить, почему алгоритм решил, что данный документ релевантен, кроме как на уровне близости эмбедингов). Кроме того, глубокие модели вроде BERT требуют значительных вычислительных ресурсов и объёмов обучающих данных [Devlin, 2019]. Они могут пропускать важные логические связи: если ответ на запрос требует многошагового вывода или соединения фактов, типичная векторная модель не сможет явно выполнить такое логическое заключение.

Графовые модели требуют поддержания сложной инфраструктуры (графовые базы данных, механизмы актуализации знаний) и зачастую специфичны для предметной области.

Современные исследования всё чаще стремятся объединить сильные стороны графов и нейросетевых моделей [Scarselli, 2009, Pan, 2024], чтобы добиться одновременно глубокого понимания текста и учёта явных знаний, поэтому необходимо проанализировать гибридные подходы в построении современных моделей.

Осознавая достоинства и недостатки как векторных, так и графовых подходов, разрабатываются гибридные методы поиска, сочетающие оба представления [Raj, 2025]. В таких системах компоненты на основе графов знаний могут обеспечивать дополнительный контекст и правила для интерпретации запроса или фильтрации результатов, тогда как нейросетевые модели векторного поиска отвечают за обобщающую способность и работу с неструктурированным текстом. Однако на современном этапе гибридные методы находятся в стадии становления. Кроме того, пока не выработано единых стандартов, как именно комбинировать оценки релевантности из векторной и графовой частей: исследуются различные схемы ранжирования, от простого сложения весов до сложных моделей на основе машинного обучения.

Сравнительный анализ

В табл. 1 приведены результаты сравнительного анализа моделей поиска, а в табл. 2 сведены их достоинства и недостатки.

Таблица 1
Table 1

Сравнение поисковых моделей
Comparison of search models

Критерий сравнения	Тип модели		
	векторная	графовая	гибридная
1	2	3	4
Основной объект представления	Текст (слова, предложения, документы)	Сущности и связи (узлы и рёбра графа)	Текст + сущности + связи
Тип знаний	Статистический, латентно-семантический	Явный, структурированный (факты, онтологии)	Смешанный: явный + латентный
Интерпретируемость результатов	Низкая (сложно объяснить, почему документ релевантен)	Высокая (путь в графе служит объяснением)	Средняя / Высокая (зависит от реализации)
Способность к логическому выводу	Отсутствует	Присутствует (через графовые запросы и reasoning)	Частичная (за счёт графового компонента)
Масштабируемость	Высокая (векторные БД, ANN-индексы)	Низкая / Средняя (зависит от размера графа)	Средняя (сложность интеграции)
Работа с неструктурированным текстом	Отличная (изначально ориентирована на текст)	Слабая (требуется предварительное извлечение)	Хорошая (текст обрабатывается векторной частью)
Зависимость от внешних знаний	Низкая (обучается на корпусе текстов)	Высокая (требует актуального графа знаний)	Высокая (качество графа критично)
Вычислительная сложность	Средняя / Высокая (для глубоких моделей)	Высокая (особенно для многошаговых запросов)	Очень высокая (комбинирование подходов)

Окончание табл. 1
 End of Table 1

1	2	3	4
Актуальность информации	Зависит от обучающих данных	Зависит от частоты обновления графа	Потенциально высокая (сочетание источников)
Примеры / реализации	TF-IDF, BM25, word2vec, BERT, DPR	PageRank, Wikidata, Neo4j, R-GCN	GraphRAG, GraphEmbed, Neo4j + LLM
Область применения	Поиск по текстам, семантическое ранжирование	Ответы на фактологические вопросы, экспертные системы	Сложные вопросно-ответные системы, аналитика

Таблица 2
 Table 2

Достоинства и недостатки поисковых моделей
 Advantages and disadvantages of search models

Тип модели	Достоинства	Недостатки
1	2	3
Векторная	– Высокая масштабируемость и скорость обработки больших объемов текста благодаря эффективным индексам (ANN). – Автоматическое определение семантического сходства без необходимости ручного построения онтологии. – Возможность улавливать контекстные значения слов (особенно в моделях типа BERT), решающая проблему полисемии. – Развитая инфраструктура: векторные базы данных (Pinecone, Weaviate), библиотеки (FAISS, ScaNN), интеграция с поисковыми системами. – Хорошая обобщающая способность на больших данных.	– Низкая интерпретируемость: сложно объяснить, почему тот или иной документ считается релевантным («чёрный ящик»). – Игнорирует явные логические связи и структурированные знания (например, отношения между сущностями). – Высокая вычислительная стоимость обучения и вывода глубоких моделей (трансформеров). – Неспособность выполнять многошаговые рассуждения и логический вывод. – Зависимость от качества и объёма обучающих данных (может вносить искажения).
Графовая	– Явное представление знаний и связей между сущностями (онтологии, триплеты). – Высокая интерпретируемость: результат поиска может быть объяснен путем в графе. – Возможность выполнения сложных структурированных запросов (SPARQL) и логического вывода (рассуждений). – Точность при работе с фактическими вопросами в узких предметных областях.	– Низкая масштабируемость при работе с графами, содержащими миллиарды узлов и связей. – Высокая трудоемкость построения, поддержки и обновления графов знаний. – Неполнота данных: граф не может содержать все факты, особенно новые или узкоспециализированные. – Сложность обработки неструктурированного текста: требует этапов распознавания именованных сущностей (NER) и извлечения связей. – Динамичность данных требует постоянного обновления, что часто невозможно в режиме реального времени.

Окончание табл. 2
End of Table 2

1	2	3
Гибридная	– Сочетание семантической глубины векторных моделей и структурной точности графов. – Повышение точности и релевантности поиска за счет учета как текстовых, так и явных связей. – Улучшенная интерпретируемость: граф может служить объяснением результатов векторного поиска. – Сокращение «галлюцинаций» в генеративных моделях (LLM) за счет проверки фактов на основе графа знаний. – Возможность обработки сложных запросов, требующих объединения информации из разных источников.	– Высокая архитектурная и имплементационная сложность: требует интеграции разрозненных компонентов. – Отсутствие устоявшихся стандартов для объединения оценок релевантности (векторных и графовых). – Критическая зависимость от качества и полноты исходного графа знаний. – Риск ухудшения качества при наличии шума или ошибок в извлечении сущностей. – Дополнительные затраты на синхронизацию и поддержание релевантности данных. – Находится на ранних стадиях разработки, мало проверенных в отрасли решений.

Проведённый анализ выявил не только технические, но и концептуальные различия между тремя подходами к информационному поиску. Для более глубокого понимания их природы, ограничений и областей применимости следует рассмотреть следующие аспекты, выходящие за рамки простого перечисления достоинств и недостатков.

Первое важное различие между моделями заключается в том, что каждая из них считает «пониманием» смысла (природа «понимания» информации: статистика vs. логика vs. синтез). В частности, векторные модели работают со статистической семантикой. Их «понимание» текста основано на гипотезе распределенной семантики: значение слова определяется его контекстным окружением. Модели, такие как word2vec или BERT, не знают, что «Париж» – это столица Франции в логическом смысле. Они «знают», что слова «Париж», «Франция», «столица», «Лувр» и «Эйфелева башня» часто встречаются в схожих контекстах, и поэтому их векторы пространственно близки. Это мощный, но статистический, вероятностный метод моделирования значения. Он эффективен, пока язык ведет себя предсказуемо, но терпит неудачу в ситуациях, требующих точного логического вывода. Графовые модели реализуют логико-символическое представление знаний. Здесь «понимание» – это формальный факт, выраженный в виде тройки (субъект-предикат-объект) или онтологического отношения. Граф знаний буквально знает, что «Париж» имеет отношение «является столицей» к «Франции». Это знание является явным, дискретным и может быть использовано для логического вывода (например, если «Париж – столица Франции» и «Франция – член ЕС», то можно сделать определенные выводы). Это знание является точным, но негибким и требует предварительной формализации. Гибридные модели стремятся синтезировать статистический и логический подходы. Идея заключается в использовании сильных сторон обоих миров: гибкости и контекстной чувствительности векторных представлений для понимания текста запросов и документов, одновременно используя точные формальные знания графов для проверки фактов, установления неочевидных связей и обеспечения объяснимости. Это попытка создать понимание, более близкое к человеческому, которое является одновременно контекстуальным и логическим.

Второе важное различие между моделями заключается в способности моделей объяснять, почему она привела к тому или иному результату, имеет решающее значение для многих областей применения (медицина, юриспруденция, наука). Векторный поиск выдает результаты, но не предоставляет удобочитаемого объяснения. Максимум, что можно сказать: «Этот документ релевантен, потому что косинусное сходство между его векторным

представлением и векторным представлением запроса составляет 0,87». Это математическое объяснение, которое ничего не говорит пользователю о содержательной связи. Для пользователя это «черный ящик»: система просто «чувствует», что текст релевантен. Поиск по графу предлагает принципиально иной тип объяснимости – структурную объяснимость. Ответ на вопрос: «Почему этот документ (или сущность) релевантен?» может быть представлен в виде пути в графе. Например: «Ваш запрос содержит сущность «Илон Маск». В графе знаний эта сущность связана отношением «основатель» компаний Tesla, а документ, который мы вам показываем, содержит информацию о последних моделях Tesla. Этот путь удобочитаем и может быть визуализирован. Объяснение здесь – это трассировка связей. Гибридный подход может использовать граф для аннотирования и обогащения результатов векторного поиска. Векторная часть находит набор семантически связанных документов. Затем на эти результаты «накладывается» граф знаний, выделяющий связи между сущностями в найденных документах и сущностями в запросе. Пользователь видит не просто список ссылок, а структурированное резюме: «В этих документах упоминаются следующие лица, связанные с темой вашего запроса...» «Были обнаружены следующие связи между упомянутыми организациями...» Таким образом, объяснение строится как постобработка результатов векторного поиска с использованием графа.

Следующее важное различие между моделями заключается в адаптивности к новым знаниям и динамике изменений. Векторные модели статистически адаптируются к новым темам. Для того чтобы модель «научилась» понимать новое понятие или событие, ей необходимо предоставить большое количество текстов, содержащих это понятие в различных контекстах. Обучение или тонкая настройка больших моделей требует времени и вычислительных ресурсов. Однако, если новый термин уже встречался в текстах, эмбединги, сгенерированные на лету (например, с помощью BERT), могут более или менее адекватно представлять его даже без переобучения, при условии, что они состоят из известных подслов. Адаптация происходит путем изменения статистической картины мира в данных. Графовые модели адаптируются путем прямого добавления фактов. Появилась новая компания? Добавьте узел. У нее появился новый продукт? Добавьте ссылку. Это можно сделать мгновенно, вручную или автоматически, но это требует либо вмешательства человека (онтолога), либо точного алгоритма извлечения информации. Однако сам граф не «понимает» контекст этого нового факта, он просто хранит его. Адаптация здесь – это административное или алгоритмическое действие по изменению структуры данных. Гибридные модели потенциально могут реализовать наилучший сценарий: компонент графа может быстро обновляться новыми фактами (например, из новостной ленты), а компонент вектора может использоваться для поиска неструктурированного текста, связанного с этими новыми фактами. Например, если в граф добавляется факт о слиянии двух корпораций, гибридная система может немедленно найти свежие аналитические статьи об этом событии, даже если сам граф не содержит этих статей. Граф обеспечивает операционную структуру, а векторы – глубину контента.

Отметим еще одно важное различие, которое состоит в том, что модели имеют различную устойчивость к неопределённости и шуму. Векторные модели по своей природе устойчивы к шуму. Если текст содержит опечатку или использует нестандартное слово, контекстное встраивание может смягчить эту проблему. Они работают в непрерывном пространстве, где «почти правильное» слово всё ещё будет близко к «правильному». Они также хорошо справляются с синонимией и перефразированием. Графовые модели чрезвычайно чувствительны к точности. Если отношение неверно указано в графе, это приведёт к неверным логическим выводам. Они работают в дискретном пространстве точных фактов. Они не обрабатывают «почти», для них факт либо существует, либо нет. Синонимия для них является отдельной проблемой (например, необходимо явно указать, что «аэропорт Шереметьево» и «SVO» – это одно и то же). Гибридные модели пытаются найти баланс. Векторная часть может «подсказывать» графовой части, какие сущности, вероятно, присутствуют в тексте (даже если они названы неканонически), а графовая часть может фильтровать зашумлённые результаты векторного поиска, удаляя те, которые противоречат известным фактам. Если в графе есть ошибки, гибридный поиск тоже начинает ошибаться. Причем ошибка может выглядеть как

точный факт, то есть то, что называют галлюцинациями модели. Это опаснее, чем обычная неточность векторной модели. Векторная модель чаще дает вероятностно близкий результат. А графовая часть может закрепить неверную связь как правильную.

Из этого следует, что ни один подход нельзя считать универсальным. Векторные модели хорошо работают с текстом. Они находят смысловую близость и подходят для больших массивов данных. Но они плохо показывают явные связи между объектами. Графовые модели, наоборот, хорошо описывают структуру знаний. Они позволяют объяснить результат через связи и отношения. Но им сложнее работать с большими объемами данных и сырым текстом. Гибридные модели выглядят наиболее перспективно, потому что соединяют оба подхода. Но пока это направление ещё развивается. Для него нужны более понятные правила построения, оценки качества и практического применения. Результаты анализа позволяют сравнить различные подходы и могут служить основой для выбора архитектуры будущей поисковой системы в зависимости от конкретных требований задачи (точность, масштабируемость и интерпретируемость).

Рекомендации по практическому совершенствованию и развитию гибридных моделей информационного поиска

Целесообразно разработать подходы к объединению пространств представлений для гетерогенных данных. Векторные текстовые встраивания и графовые представления сущностей работают в разных семантических измерениях, что затрудняет их совместное использование. Один из возможных путей развития – совместное обучение векторной и графовой части модели. Сейчас они часто работают отдельно. Векторная часть ищет близкие по смыслу документы. Графовая часть хранит связи между сущностями. Лучше, если эти два механизма будут настраиваться под одну задачу поиска. Тогда близость между запросом и документом будет определяться не только похожими словами или эмбедингами. Она также будет учитывать связи в графе знаний.

Еще одна важная задача – лучше извлекать знания из обычного текста. Для построения графа нужно находить сущности и связи между ними. Современные нейросетевые методы справляются с этим лучше прежних подходов. Но ошибки всё равно остаются. Неверно найденная сущность или неправильная связь могут попасть в граф. После этого они будут влиять на результаты поиска.

Поэтому новые факты перед добавлением в граф нужно проверять. Для этого можно сравнивать данные из нескольких источников, составив из части информации ложные и истинные триплеты. А после сравнивать их на нестыковки. Если факты конфликтуют друг с другом, то эти участки информации нужно вручную обработать, чтобы человек нашел нестыковки.

Также нужно улучшать ранжирование в гибридных системах. Нельзя всегда заранее задавать один и тот же вес для векторной и графовой части. В одних запросах важнее смысловая близость текста. В других важнее точная связь между сущностями. Поэтому лучше использовать обучаемые модели ранжирования. Они смогут сами определять, какая часть системы важнее для конкретного запроса. Например, для фактических запросов, требующих точного знания связей между сущностями, приоритет может быть отдан графовому компоненту, в то время как для тематических или аналитических запросов более важны векторные представления, отражающие семантическое сходство текстов.

С точки зрения архитектурной реализации гибридных систем, перспективным представляется разработка многоуровневых конвейеров обработки запросов. На первом уровне целесообразно использовать быстрые методы для грубого отбора кандидатов, например, на основе обратного индекса или приблизительного поиска ближайшего соседа в векторном пространстве. На втором уровне полученный набор кандидатов обогащается информацией из графа знаний: для каждого документа определяются связанные сущности, устанавливаются пути графа, соединяющие запрос с документом, и вычисляются метрики

структурного сходства. На третьем уровне выполняется окончательное ранжирование с использованием составной модели, учитывающей как текстовую релевантность, так и выявленные связи в графе. Такая иерархическая архитектура позволяет сочетать высокую производительность на этапе отбора с семантической глубиной на этапе ранжирования.

Важным аспектом разработки гибридных моделей является повышение их интерпретируемости. Рекомендуются разработать методы визуализации и объяснения результатов поиска, которые отображают не только ссылки на документы, но и фрагменты графа знаний, подтверждающие релевантность. Например, система может показывать путь графа, соединяющий сущности запроса с сущностями, упомянутыми в документе, или выделять ключевые связи, которые привели к высокой оценке релевантности. Такой подход делает результаты поиска понятнее для пользователя. Он видит не только найденный документ, но и связь, по которой система пришла к этому результату. Кроме того, пользователь может уточнять запрос через элементы графа знаний. Например, он может перейти от одной сущности к связанной с ней сущности и сузить область поиска. На практике гибридные системы лучше строить не с нуля, а на основе уже готовых инструментов. Для графовой части можно использовать графовые базы данных. Для текстовой части – векторные хранилища. Такие решения уже позволяют хранить эмбединги, искать похожие документы и работать со связями между сущностями. Но главная задача здесь – правильно связать эти части между собой. Нужно заранее продумать, как будет выполняться общий запрос. Например, когда сначала используется векторный поиск, а затем результат уточняется через связи в графе. Или наоборот, когда граф задает область поиска, а векторная модель выбирает наиболее близкие документы. В таких системах граф знаний выступает в качестве внешнего хранилища проверенных фактов, к которому языковая модель обращается для проверки своих гипотез и предоставления контекста. Рекомендуются разработать методы эффективного взаимодействия между языковыми моделями и графами знаний, включая преобразование неформальных запросов пользователей в формальные запросы к графу и обратную интерпретацию результатов поиска по графу на естественном языке.

Наконец, ключевой областью развития гибридных моделей является создание эталонных наборов данных и методов оценки, учитывающих специфику гибридного поиска. Традиционные метрики, такие как точность и полнота, хотя и остаются важными, не в полной мере отражают способность системы использовать структурные знания. Рекомендуются разработать комплексные методологии тестирования, включающие сценарии, требующие многоэтапного рассуждения, установления неочевидных связей между сущностями и проверки фактической достоверности результатов. Создание общедоступных эталонных наборов данных для гибридного поиска позволит объективно сравнивать различные подходы и стимулировать дальнейшие исследования в этой области.

Заключение

В работе проведён анализ основных подходов к представлению и поиску информации. Векторные модели (от классического TF-IDF до современных нейросетевых эмбедингов BERT) обеспечивают эффективную обработку текстовых данных и способны выявлять семантическую близость, однако не отражают явно структуру знаний и причинно-логические связи. Графовые модели (графы знаний, онтологии, а также методы на основе графовых нейронных сетей) позволяют учесть отношения между сущностями и осуществлять логический вывод, но страдают от проблем масштабируемости и неполноты, что затрудняет их прямое применение на больших неструктурированных коллекциях. Рассмотрены современные гибридные подходы, объединяющие графы и эмбединги: они показывают перспективность (улучшение точности и интерпретируемости поиска), но находятся на ранней стадии развития и пока не решают всех проблем в полной мере.

Таким образом, был выявлен существенный научный пробел: отсутствует универсальная методология, позволяющая в рамках единого алгоритма поиска задействовать и богатство

графовых знаний, и гибкость векторных представлений. Поэтому возникает необходимость глубокой разработки нового подхода в теории поиска информации, способного динамически интегрировать результаты на основе семантического контекста запроса, сохраняя при этом масштабируемость.

Список литературы

- Абрамович Р.К., Добрынин В.Ю., Платонов А.В. 2025. Объединение глубоких моделей и разреженных представлений в информационном поиске: обзор и анализ современных подходов. *Информационные и математические технологии в науке и управлении*, 2(38): 5–17. DOI:10.25729/ESI.2025.38.2.001.
- Настасенко С.А., Савотченко С.Е. 2025. Гибридный метод поиска в текстовых данных с использованием графов и эмбедингов. *Новые идеи в науках о Земле*. Москва: Российский государственный геологоразведочный университет им. Серго Орджоникидзе, 148–151.
- Настасенко С.А., Савотченко С.Е. 2025. Интеграция графовых структур в эмбединги векторных баз данных для повышения семантического поиска. *Донецкие чтения – 2025: образование, наука, инновации, культура и вызовы современности*. Донецк: Изд-во ДонГУ, 1: 222–224. EDN: NUFFAH
- Chang Y., Wang, X., Wang J., Wu Y., Yang L., Zhu K., Xie X. 2024. A survey on evaluation of large language models. *ACM transactions on intelligent systems and technology*, 15(3), 1–45.
- Chao Li, Hao Xu, Kun He, (2024). Meta-multigraph search: Rethinking meta-structure on heterogeneous information networks, *Knowledge-Based Systems*, 289:111524. <https://doi.org/10.1016/j.knsys.2024.111524>.
- Chatterjee S., Dalton J. 2025. QDER: Query-Specific Document and Entity Representations for Multi-Vector Document Re-Ranking. *Proceedings of the 48th International ACM SIGIR Conference on Research and Development in Information Retrieval*: 2255–2265. <https://doi.org/10.1145/3726302.3730065>.
- Devlin J., Chang M.-W., Lee, K., Toutanova K. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 1: 4171–4186. <https://doi.org/10.18653/v1/N19-1423>.
- Kruhlov, V., Bobos, O., Hnylianska, O., Rossikhin, V., & Kolomiets, Y. 2024. The Role of Using Artificial Intelligence for Improving the Public Service Provision and Fraud Prevention. *Pakistan Journal of Criminology*, 16(2): 913–928.
- Kyurkchiev P., Iliev A., Kyurkchiev N. 2026. Semantic Search for System Dynamics Models Using Vector Embeddings in a Cloud Microservices Environment. *Future Internet*, 18(2):86. <https://doi.org/10.3390/fi18020086>.
- Pan Y., Wang L. 2024. Retrieval-Augmented Generation for Code Summarization via Hybrid GNN. *IEEE Trans. Softw. Eng.* 50: 1234–1250.
- Papageorgiou G, Sarlis V, Maragoudakis M, Tjortjis C. 2025. Hybrid Multi-Agent GraphRAG for E-Government: Towards a Trustworthy AI Assistant. *Applied Sciences*. 15(11): 6315. <https://doi.org/10.3390/app15116315>
- Papageorgiou G, Sarlis V, Maragoudakis M, Tjortjis C. 2025. A Multimodal Framework Embedding Retrieval-Augmented Generation with MLLMs for Eurobarometer Data. *AI*. 6(3): 50. <https://doi.org/10.3390/ai6030050>
- Raj M., Mishra, N. 2025. Exploring new Approaches for Information Retrieval through Natural Language Processing, *arXiv e-prints*, Art. no. *arXiv:2505.02199*, doi:10.48550/arXiv.2505.02199.
- Sadykova T., Sinchev B., Young I. C., Auyezova A. 2025. The application of vector space models in intelligent information retrieval systems. *Academic Scientific Journal of Computer Science*, 355(3), 160–175. <https://doi.org/10.32014/2025.2518-1726.370>
- Sarmah B., Hall B., Rao R., Patel S., Pasquali S., Mehta D. 2024. HybridRAG: Integrating Knowledge Graphs and Vector Retrieval Augmented Generation for Efficient Information Extraction. *arXiv preprint arXiv:2408.04948*.
- Scarselli F., Gori M., Tsoi A.C., Hagenbuchner M., Monfardini G. 2009. The Graph Neural Network Model. *IEEE Trans. Neural Netw.* 20, 61–80. <https://doi.org/10.1109/TNN.2008.2005605>
- Zhu X., Guo X., Cao S., Li S., Gong J. 2024. StructuGraphRAG: Structured document-informed knowledge graphs for retrieval-augmented generation. In *Proceedings of the AAAI symposium series*. 4(1): 242–251.

References

- Abramovich R.K., Dobrynin V.YU., Platonov A.V. 2025. Ob'edinenie glubokih modelej i razrezhennyh predstavlenij v informacionnom poiske: obzor i analiz sovremennyh podhodov [Combining deep models and sparse representations in information retrieval: A review and analysis of modern approaches]. *Informacionnye i matematicheskie tekhnologii v nauke i upravlenii*, 2(38) 5–17. DOI:10.25729/ESI.2025.38.2.001.
- Nastasenکو S.A., Savotchenko S.E. 2025. Gibridnyj metod poiska v tekstovyh dannyh s ispol'zovaniem grafov i embeddingov [A Hybrid Method for Searching Text Data Using Graphs and Embeddings]. *Novye idei v naukah o Zemle*. Moskva: Rossijskij gosudarstvennyj geologorazvedochnyj universitet im. Sergo Ordzhonikidze, 148–151.
- Nastasenکو S.A., Savotchenko S.E. 2025. Integraciya grafovyyh struktur v embenddingi vektornyh baz dannyh dlya povysheniya semanticheskogo poiska [Integrating graph structures into vector database embeddings to improve semantic search]. *Doneckie chteniya – 2025: obrazovanie, nauka, innovacii, kul'tura i vyzovy sovremennosti*. Doneck: Izd-vo DonGU, 1: 222–224. EDN: NUFFAH
- Chang Y., Wang X., Wang J., Wu Y., Yang L., Zhu K., Xie X. 2024. A survey on evaluation of large language models. *ACM transactions on intelligent systems and technology*, 15(3), 1–45.
- Chao Li, Hao Xu, Kun He, (2024). Meta-multigraph search: Rethinking meta-structure on heterogeneous information networks, *Knowledge-Based Systems*, 289: 111524. <https://doi.org/10.1016/j.knosys.2024.111524>.
- Chatterjee, S., Dalton, J. 2025. QDER: Query-Specific Document and Entity Representations for Multi-Vector Document Re-Ranking. *Proceedings of the 48th International ACM SIGIR Conference on Research and Development in Information Retrieval*: 2255–2265. <https://doi.org/10.1145/3726302.3730065>.
- Devlin J., Chang M.-W., Lee, K., Toutanova K. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 1: 4171–4186. <https://doi.org/10.18653/v1/N19-1423>.
- Kruhlov, V., Bobos, O., Hnylianska, O., Rossikhin, V., & Kolomiets, Y. 2024. The Role of Using Artificial Intelligence for Improving the Public Service Provision and Fraud Prevention. *Pakistan Journal of Criminology*, 16(2): 913–928.
- Kyurkchiev P., Iliev A., Kyurkchiev N. 2026. Semantic Search for System Dynamics Models Using Vector Embeddings in a Cloud Microservices Environment. *Future Internet*, 18(2): 86. <https://doi.org/10.3390/fi18020086>.
- Pan Y., Wang L. 2024. Retrieval-Augmented Generation for Code Summarization via Hybrid GNN. *IEEE Trans. Softw. Eng.* 50:1234–1250.
- Papageorgiou G, Sarlis V, Maragoudakis M, Tjortjis C. 2025. Hybrid Multi-Agent GraphRAG for E-Government: Towards a Trustworthy AI Assistant. *Applied Sciences*. 15(11): 6315. <https://doi.org/10.3390/app15116315>
- Papageorgiou G, Sarlis V, Maragoudakis M, Tjortjis C. 2025. A Multimodal Framework Embedding Retrieval-Augmented Generation with MLLMs for Eurobarometer Data. *AI*. 6(3): 50. <https://doi.org/10.3390/ai6030050>
- Raj M., Mishra, N. 2025. Exploring new Approaches for Information Retrieval through Natural Language Processing, *arXiv e-prints, Art. no. arXiv:2505.02199*, doi:10.48550/arXiv.2505.02199.
- Sadykova T., Sinchev B., Young I.C., Auyezova A. 2025. The application of vector space models in intelligent information retrieval systems. *Academic Scientific Journal of Computer Science*, 355(3), 160–175. <https://doi.org/10.32014/2025.2518-1726.370>
- Sarmah B., Hall B., Rao R., Patel S., Pasquali S., Mehta D. 2024. HybridRAG: Integrating Knowledge Graphs and Vector Retrieval Augmented Generation for Efficient Information Extraction. *arXiv preprint arXiv:2408.04948*. <https://doi.org/10.48550/arXiv.2408.04948>
- Zhu X., Guo X., Cao S., Li S., Gong J. 2024. StructuGraphRAG: Structured document-informed knowledge graphs for retrieval-augmented generation. In *Proceedings of the AAAI symposium series*, 4(1): 242–251.

Конфликт интересов: о потенциальном конфликте интересов не сообщалось.

Conflict of interest: no potential conflict of interest related to this article was reported.

Поступила в редакцию 25.02.2026

Поступила после рецензирования 24.05.2026

Принята к публикации 29.05.2026

Received February 25, 2026

Revised May 24, 2026

Accepted May 29, 2026



ИНФОРМАЦИЯ ОБ АВТОРАХ

Настасенко Сергей Александрович, аспирант кафедры высшей математики и физики, Российский государственный геологоразведочный университет имени Серго Орджоникидзе, г. Москва, Россия

Савотченко Сергей Евгеньевич, доктор физико-математических наук, профессор кафедры высшей математики и физики, Российский государственный геологоразведочный университет имени Серго Орджоникидзе, г. Москва, Россия

INFORMATION ABOUT THE AUTHORS

Sergey A. Nastasenko, Postgraduate Student of the Department of Higher Mathematics and Physics, Sergo Ordzhonikidze Russian State University for Geological Prospecting, Moscow, Russia

Sergey E. Savotchenko, Doctor of Physical and Mathematical Sciences, Professor of the Department of Higher Mathematics and Physics, Sergo Ordzhonikidze Russian State University for Geological Prospecting, Moscow, Russia