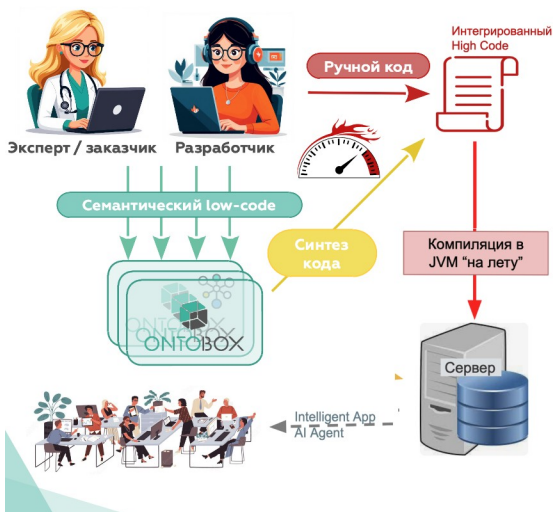




Ontobox + LLM: стратегия интеграции

1 Ontobox – платформа семантического low-code.....	1
2 Ключевые преимущества интеграции семантических моделей Ontobox и LLM.....	4
2.2 Этап эксплуатации: рой AI-агентов.....	5
3 Ontobox + LLM как новая архитектура RAG.....	7

1 Ontobox – платформа семантического low-code



- **Семантическая модель:**
процессы сущности и операции на языке бизнеса
- Эксперты-предметники как **разработчики**: понимание структуры приложения, быстрая обратная связь
- **Автоматическая генерация** до 90% системы:
бэкенд-логики, API, UI
- Системы **любых масштабов** – модели как микросервисы
- Резкое **ускорение разработки**: 1-2 специалиста создают управленческие системы и рой AI-агентов

Ontobox – это **инновационная платформа для быстрого создания управленческих систем, цифровых двойников, интеллектуальных приложений и ИИ-агентов**. Основана на принципах семантического моделирования. Сочетает **low-code подход с технологиями больших языковых моделей**.

Ключевые характеристики Ontobox:

Прозрачные семантические модели вместо ручного кодирования. Разработчик описывает предметную область в виде **объектной онтологии** (семантической модели) – сущности, связи, бизнес-правила. На основе этой модели Ontobox автоматически генерирует до **90% серверного кода** системы. Это радикально ускоряет разработку и снижает зависимость от больших команд программистов. Всего **1–2 разработчика** с помощью Ontobox могут создать и сопровождать систему масштаба корпорации, что традиционно требовало бы целого отдела. При этом семантическая модель читаема для бизнес-экспертов, они могут напрямую участвовать в её создании и верификации ("expert-in-the-loop"), реализуя возможность глубокого проникновения экспертных компетенций в процесс разработки и уменьшая риск недопонимания требований. Прозрачность моделей, которые являются ядром разрабатываемых управленческих систем, является основой для высокоуровневого

Ontobox устраняет проблемы классических подходов:

- *Долгая и дорогая разработка на чистом коде:* платформа сокращает **time-to-market** за счёт генерации кода из моделей, позволяя в разы быстрее получить рабочий прототип или продукт.
- *Ограниченность типичных low-code платформ:* большинство low-code нацелены на простые сценарии и являются чёрными ящиками. Ontobox же поддерживает **проекты любого масштаба** – от MVP для стартапа до **цифрового двойника** холдинга – и генерирует открытый прозрачный код. Системы легко расширять и модифицировать за счёт эволюции модели, без «тупиков», свойственных закрытым платформам.
- *Интеграция ИИ без потери контроля:* подключение LLM обычно связано с риском неконтролируемых выводов модели («галлюцинаций»). В Ontobox же большая

языковая модель работает в управляемом контексте – знания заложены в онтологию. Семантическая модель служит единым “источником правды” для ИИ-агентов, что повышает надёжность их работы. Ontobox по сути связывает символический ИИ (онтологии) и генеративный ИИ, позволяя получать преимущества обоих подходов.

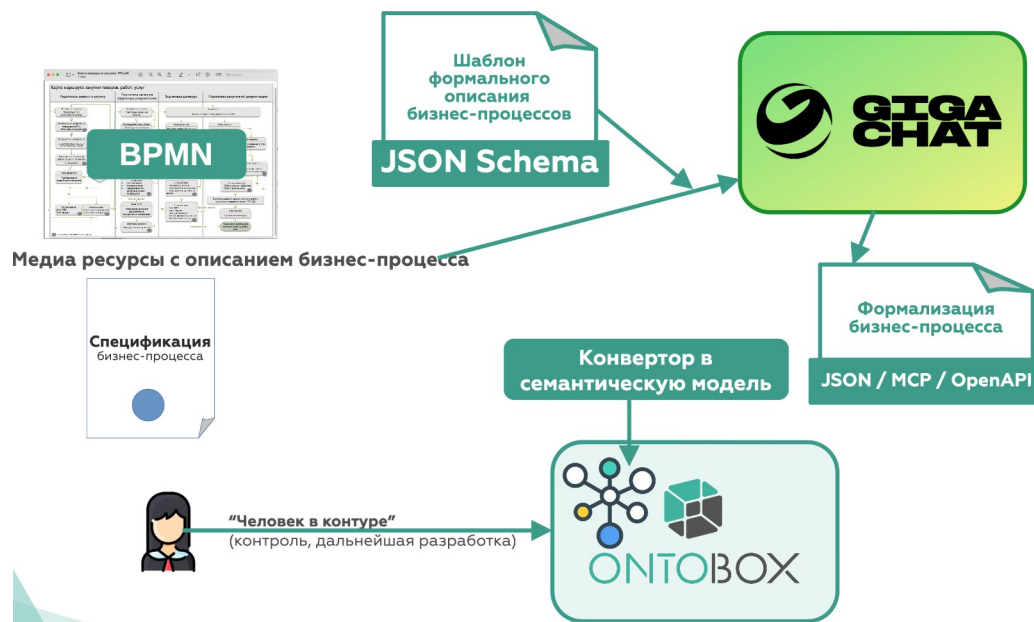
- **Импортонезависимая, промышленная реализация:** Платформа может быть развёрнута on-premise, работает с отечественными технологиями. Это важно для корпоративных клиентов: Ontobox обеспечивает полный **контроль над данными** и отсутствие зависимости от иностранных сервисов, что актуально для российских реалий. Решения на Ontobox легко интегрируются в существующую ИТ-инфраструктуру (БД, микросервисы, очереди и др.), поддерживают масштабирование и отказоустойчивость на уровне enterprise.

Ontobox представляет собой технологичную платформу, сочетающую **семантическое моделирование, low-code, базирующийся на моделировании, и LLM** и ориентирован на новые способы применения больших языковых моделей в бизнесе. Платформа **сокращает разрыв между бизнесом и ИТ**, позволяя быстрее и дешевле создавать умные корпоративные системы.

2 Ключевые преимущества интеграции семантических моделей Ontobox и LLM

Интеграция платформы Ontobox с LLM даёт синергетический эффект на **двух основных этапах жизненного цикла** управленческих систем – на стадии разработки и на стадии эксплуатации. Ниже выделены главные преимущества каждого из них:

2.1 Этап разработки: Vibe-Modeling



Ускоренная разработка (vibe-modeling): при создании новой системы Ontobox использует LLM для автоматической генерации семантических моделей из неструктурированных медийных материалов. Этот процесс мы называем **vibe-modeling**. Он позволяет разработчику значительно сэкономить время на первоначальном проектировании онтологии. Например, достаточно предоставить LLM описания процессов на естественном языке, транскрипты обсуждений, требования или диаграммы процесса – и модель сгенерируется автоматически.

Как это работает: разработчики задают LLM мультимодальные материалы (текст, аудио расшифровки, схемы) и **JSON-схему** домена (как вариант – могут быть заданы MCP-описания или описания в формате OpenAPI), по которой нужно представить знания. LLM анализирует информацию и выдаёт **JSON-модель бизнес-процесса**, удовлетворяющую заданной схеме. Затем Ontobox конвертирует этот JSON в **черновик объектной онтологии** в своей среде. По сути, LLM выступает в роли интеллектуального ассистента, который понимает текст и сразу формирует формальную модель. На практике это позволяет **автоматизировать от 50% до 100% работы** по созданию модели бэкенда. Разработка ускоряется в ~10 раз, поскольку LLM берёт на себя рутинный анализ текстов и черновое проектирование структур данных. Разработчику остаётся лишь откорректировать и доработать модель, вместо того чтобы строить её с нуля. Этот

сочетанный подход (генерация модели LLM + семантический low-code) даёт существенное преимущество в скорости и гибкости разработки.

2.2 Этап эксплуатации: рой AI-агентов



Интеллектуальная эксплуатация (рой контекстных AI-агентов): после запуска созданной на Ontobox системы интеграция с LLM позволяет вывести её возможности на новый уровень за счёт **набора узкоспециализированных AI-агентов**, действующих вокруг основной системы. В разработанной на Ontobox управленческой/корпоративной системе заложена формальная логически корректная модель предметной области (онтология), и она используется как **единый контекст** для всех взаимодействующих с ней ИИ-агентов.

Как это работает:

1. Для каждой интеллектуальной задачи, требующей взаимодействия с LLM, в рамках материнской управленческой системы создаётся свой “узкопрофильный” агент

(например, помощник врача, анализатор документов, консультант по данным и т.п.).

2. Каждый агент через API получает из Ontobox-модели релевантный **контекст (структурированные данные, бизнес-правила, словари)** и формирует запрос к LLM в этом контексте.
3. **LLM, опираясь на связную картину мира, предоставленную целостной моделью Ontobox, выдаёт целевой результат.**

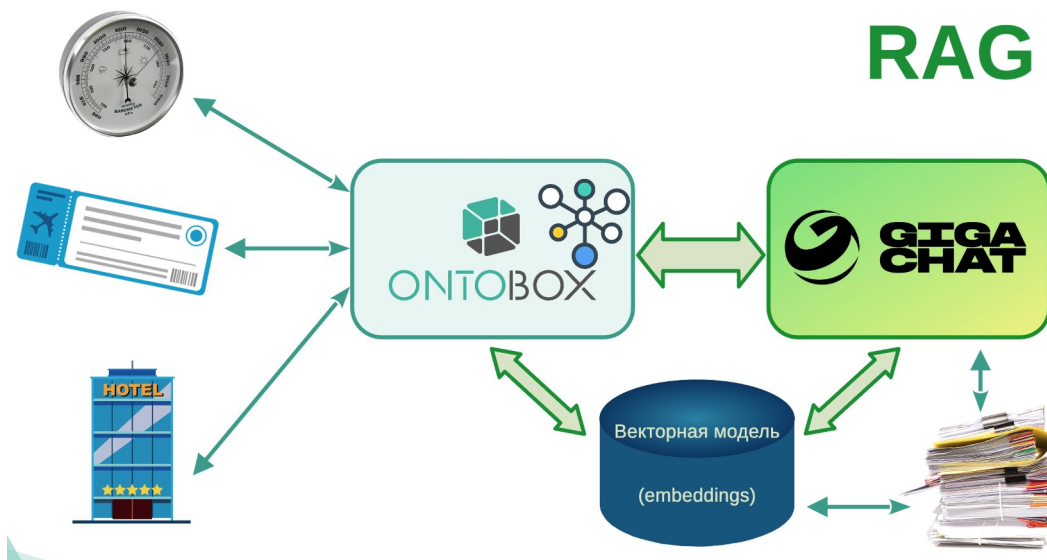
Благодаря такому подходу:

- ИИ-агенты работают согласованно, **говоря «на одном языке»** с основной системой. Материнская система на Ontobox служит как бы **«цифровым мозгом»** для роя агентов, предотвращая противоречия и расхождения.
- **Снижается риск ошибок и галлюцинаций** у модели. Поскольку на вход LLM получают не произвольный текст, а *структурированный контекст с фактами и ограничениями*, ответы более достоверны.
- Можно масштабировать интеллект системы горизонтально: добавлять новых агентов под разные нужды, не ломая общую архитектуру. **Единая семантическая среда** обеспечивает целостность знаний при большом количестве микросервисных AI-компонентов. Такой **рой ИИ** вокруг Ontobox-системы увеличивает её функциональность *без переписывания кода*: достаточно обучить агента работать с контекстом онтологии и нужными запросами в LLM.

В итоге интеграция Ontobox + LLM создаёт **замкнутый цикл ИИ в корпоративной системе**: LLM ускоряет **создание** системы (помогая спроектировать её мозг – онтологию), а затем LLM же на этапе **эксплуатации** становится частью этой системы в виде множества агентов, действующих в строгом контексте. Такой подход сочетает сильные стороны человека и ИИ на всех этапах и формирует конкурентное преимущество нашего продукта.

Такая интеграция позволяет раскрывать потенциал LLM в рамках стратегии деплоймента больших языковых моделей в среды индустриального уровня.

3 Ontobox + LLM как новая архитектура RAG



Современные архитектуры RAG становятся важнейшей компонентой интеллектуальных систем, объединяя генеративные модели с внешними источниками знаний. Однако большинство существующих RAG-систем страдает от семантической раздробленности, отсутствия общего контекстного слоя и высокой вероятности «галлюцинаций» со стороны языковой модели. Гибрид Ontobox + LLM позволяет переосмыслить концепцию RAG как управляемую, интерпретируемую и расширяемую архитектуру.

1. Семантическая модель как ядро контекста. Ontobox предоставляет в качестве основы RAG объектно-ориентированную онтологию, которая описывает ключевые сущности, связи и бизнес-правила предметной области. Это не просто справочник терминов — это полная декларативная модель логики системы, которую LLM использует как контекстную опору для reasoning и генерации. Такая модель действует как «когнитивное воображение» системы, формируя чёткое «видение» предметной области в целом.

2. Семантическая шина и согласованность компонентов. Благодаря своей архитектуре Ontobox формирует единую «семантическую шину» — общий слой координации всех компонентов RAG. Разрозненные модули (поиск, агрегация, генерация, внешние инструменты) действуют не автономно, а внутри согласованного семантического пространства и на едином языке этого пространства. Это особенно важно для

масштабируемых корпоративных решений, где согласованность контекста между модулями критична для надёжности.

3. Семантика инструментов: декларативные методы в модели. Онтология Ontobox описывает не только объекты и связи, но и сигнатуры методов — транзакций, процедур, функций, доступных для исполнения. Это значит, что в рамках RAG мы получаем не только доступ к знаниям, но и знания о действиях. Для LLM это означает, что она может использовать эти методы как управляемые «tools» — например, триггерить процедуры, задавать параметры, выполнять действия внутри бизнес-контекста. Семантически связанное описание структур данных и действий в рамках одной модели повышает уровень стабильности и надёжности работы системы, избавляет от проблем двойных толкований и противоречий.

4. Семантическое посредничество для внешних источников. Вместо прямого подключения множества источников данных к LLM, гибридная архитектура строится через Ontobox как семантический посредник. Ontobox получает данные, преобразует их в структурированную форму, встраивает в онтологию, и только затем предоставляет LLM доступ через формализованный и единообразный семантический интерфейс. Это обеспечивает единообразие, валидность и защиту от контекстного шума и ошибок интерпретации. В данной архитектуре LLM взаимодействует с внешними источниками через семантический “guard” модели, которая контролирует смысловое единообразие данных, единый контекст для них, является фильтром/конвертором для неструктурированных и некорректно структурированных данных.

5. Снижение галлюцинаций и повышение надёжности. LLM, действуя внутри заранее описанного и формализованного семантического пространства Ontobox, существенно снижает риск генерации некорректных или недостоверных ответов. Модель действует в рамках целостного контекста, а не на основе обрывочных неоднородных знаний. Это делает гибридную архитектуру особенно перспективной для критически важных систем: в медицине, финансах, управлении.

6. Архитектура роя агентов с общей онтологической базой RAG на базе Ontobox поддерживает концепцию роя ИИ-агентов, каждый из которых решает

узкоспециализированные задачи (например, подбор лекарств, проверка контрактов, анализ графиков). Все агенты работают в рамках одной онтологии — это исключает конфликты, дублирование и фрагментированность логики. Мы получаем модульную, расширяемую, но при этом полностью согласованную интеллектуальную систему.

7. LLM + векторная модель как источники развития онтологий. LLM и векторные модели через семантический анализ неструктурированных мультимодальных данных могут экстрагировать из них формальные описания, используемые как строительный материал при разработке онтологий. Таким образом обеспечивается двусторонняя синергия: с одной стороны Ontobox контролирует целостность семантического пространства RAG, единообразное функционирование внешних источников данных, а с другой стороны, LLM+Embedding позволяют семантически структурировать мультимодальные данные для обогащения ими семантических моделей Ontobox, тем самым резко ускоряя разработку.

8. Метафора: воображение и интуиция Если сравнивать такую архитектуру с человеческим интеллектом, Ontobox выполняет роль когнитивного воображения — структурной основы мышления и логики, тогда как LLM — это интуиция и гибкость reasoning. Вместе они дают мощную, более качественно управляемую систему.

Гибрид Ontobox + LLM в архитектуре RAG — это шаг вперёд от фрагментированных, трудноуправляемых систем к целостным, прозрачным и масштабируемым интеллектуальным средам. Он даёт разработчикам мощные инструменты, заказчикам — доверие и надёжность, а самим системам — способность учиться и адаптироваться в рамках заданной семантики. Это не просто RAG, это RAG нового поколения.