

ПРОСТОЕ РУКОВОДСТВО ПО RAG

Поисково-дополненная генерация

Абхинав Кимоти



БОМБОРА
ИЗДАТЕЛЬСТВО

Москва

УДК 004.8
ББК 32.813
К40

Abhinav Kimothi
A SIMPLE GUIDE TO RETRIEVAL AUGMENTED GENERATION

© Publishing House EKSMO 2026. Authorized translation of the English edition © 2025 Manning Publications. This translation is published and sold by permission of Manning Publications, the owner of all rights to publish and sell the same.

Кимоти, Абхинав.
К40 Простое руководство по RAG : поисково-дополненная генерация / Абхинав Кимоти ; [перевод с английского Ю. Н. Смирновой]. — Москва : Эксмо, 2026. — 288 с. — (Manning: профессиональные книги для ИТ-специалистов).

ISBN 978-5-04-232419-2

Это исчерпывающее руководство по технологии Retrieval Augmented Generation (RAG), которая делает современные большие языковые модели (LLM) точными и надежными. Абхинав Кимоти пошагово объясняет архитектуру RAG: от создания баз знаний и конвейеров индексирования до продвинутых стратегий оптимизации и оценки. Практические примеры на Python помогут инженерам и разработчикам легко внедрить RAG в свои проекты. Отличный инструмент для освоения передового AI — от основ до экспертных решений в реальных задачах.

УДК 004.8
ББК 32.813

ISBN 978-5-04-232419-2

© Смирнова Ю.Н., перевод на русский язык, 2026
© Оформление. ООО «Издательство «Эксмо», 2026

Краткое содержание

Предисловие	11
Благодарности	13
О книге	15
ЧАСТЬ 1. ОСНОВЫ	19
Глава 1. LLM и необходимость RAG	21
Глава 2. RAG-системы и их проектирование	37
ЧАСТЬ 2. СОЗДАНИЕ RAG-СИСТЕМ	51
Глава 3. Пайплайн индексации: создание базы знаний для RAG	53
Глава 4. Пайплайн генерации: создание контекстных ответов с помощью LLM	82
Глава 5. Оценка RAG: точность, релевантность и достоверность	116
ЧАСТЬ 3. RAG В ПРОДАКШЕНЕ	153
Глава 6. Процесс развития RAG-систем: наивный, продвинутый и модульный RAG	155
Глава 7. Эволюция стека RAGOps	182
ЧАСТЬ 4. ДОПОЛНИТЕЛЬНЫЕ СООБРАЖЕНИЯ	205
Глава 8. Варианты RAG: графовый, мультимодальный, агентный и другие	207
Глава 9. Общая схема разработки RAG и направления для дальнейшего изучения	246
Указатель	278

Содержание

Предисловие	11
Благодарности	13
О книге	15
Как организована эта книга: краткий обзор	16
О коде	17
Форум на liveBook	18

ЧАСТЬ 1. ОСНОВЫ

Глава 1. LLM и необходимость RAG	21
1.1. Проклятие LLM и идея RAG	22
1.1.1. LLM не обучаются на фактах	24
1.1.2. Что такое RAG?	26
1.2. Новизна подхода RAG	29
1.2.1. Открытие RAG	29
1.2.2. Как помогает RAG?	30
1.3. Популярные сценарии использования RAG	32
1.3.1. Взаимодействие с поисковой системой	32
1.3.2. Генерация персонализированного маркетингового контента	32
1.3.3. Комментирование событий в режиме реального времени	32
1.3.4. Диалоговые агенты	33
1.3.5. Системы вопросно-ответного поиска в документах (QA-системы)	33
1.3.6. Виртуальные помощники	34
1.3.7. Исследования с использованием ИИ	34
1.3.8. Мониторинг социальных сетей и анализ тональности текстов	34
1.3.9. Генерация новостей и курирование контента	35
<i>Резюме</i>	35
Глава 2. RAG-системы и их проектирование	37
2.1. Как устроена RAG-система?	38
2.2. Проектирование RAG-систем	42
2.3. Пайплайн индексации	43

2.4. Пайплайн генерации	46
2.5. Оценка и мониторинг	47
2.6. Стек RAGOps	48
2.7. Кеширование, механизмы контроля и защиты, безопасность и другие слои	49
<i>Резюме</i>	50

ЧАСТЬ 2. СОЗДАНИЕ RAG-СИСТЕМ

Глава 3. Пайплайн индексации: создание базы знаний для RAG	53
3.1. Загрузка данных	54
3.2. Разделение данных (чанкинг)	58
3.2.1. Преимущества чанкинга	59
3.2.2. Процесс чанкинга	60
3.2.3. Методы чанкинга	60
3.2.4. Выбор стратегии чанкинга	66
3.3. Преобразование данных (эмбединги)	68
3.3.1. Что такое эмбединги?	68
3.3.2. Популярные предобученные эмбединг-модели	70
3.3.3. Сценарии использования эмбедингов	72
3.3.4. Как выбрать эмбединги?	75
3.4. Хранение (векторные базы данных)	76
3.4.1. Что такое векторные базы данных?	76
3.4.2. Типы векторных баз данных	76
3.4.3. Выбор векторной базы данных	79
<i>Резюме</i>	80
Глава 4. Пайплайн генерации: создание контекстных ответов с помощью LLM	82
4.1. Обзор пайплайна генерации	83
4.2. Поиск и извлечение	84
4.2.1. Эволюция методов поиска	85
4.2.2. Популярные ретриверы	93
4.2.3. Реализация простого ретривера	94
4.3. Аугментация	95
4.3.1. Методы промпт-инжиниринга в RAG	96
4.3.2. Создание простого промпта для аугментации	102
4.4. Генерация	104
4.4.1. Классификация LLM и их пригодность для RAG	104
4.4.2. Завершение пайплайна RAG: генерация с использованием LLM	111
<i>Резюме</i>	113

Глава 5. Оценка RAG: точность, релевантность и достоверность	116
5.1. Ключевые аспекты оценки RAG	117
5.1.1. Оценки качества	118
5.1.2. Необходимые способности	118
5.2. Метрики оценки	121
5.2.1. Метрики для оценки поиска	121
5.2.2. Специализированные метрики для RAG	129
5.3. Фреймворки	134
5.3.1. RAGAs	134
5.3.2. Автоматизированная система оценки RAG	142
5.4. Бенчмарки	143
5.4.1. RGB	143
5.5. Ограничения и лучшие практики	148
<i>Резюме</i>	150

ЧАСТЬ 3. RAG В ПРОДАКШЕНЕ

Глава 6. Процесс развития RAG-систем:	
наивный, продвинутый и модульный RAG	155
6.1. Ограничения наивного RAG	156
6.2. Техники продвинутого RAG	158
6.3. Предпоисковые методы	159
6.3.1. Оптимизация индексов	160
6.3.2. Оптимизация запросов	165
6.4. Стратегии поиска	170
6.4.1. Гибридный поиск	170
6.4.2. Итеративный поиск	170
6.4.3. Рекурсивный поиск	171
6.4.4. Адаптивный поиск	171
6.5. Постпоисковые методы	172
6.5.1. Сжатие	173
6.6. Модульный RAG	175
6.6.1. Основные модули	176
6.6.2. Новые модули	177
<i>Резюме</i>	180
Глава 7. Эволюция стека RAGOps	182
7.1. Эволюция стека RAGOps	183
7.1.1. Критичные слои	184
7.1.2. Обязательные слои	192
7.1.3. Дополнительные слои	197
7.2. Лучшие практики для продакшена	200
<i>Резюме</i>	203

ЧАСТЬ 4. ДОПОЛНИТЕЛЬНЫЕ СООБРАЖЕНИЯ

Глава 8. Варианты RAG: графовый, мультимодальный, агентный и другие	207
8.1. Что такое варианты RAG и для чего они нужны	208
8.2. Мультимодальный RAG	209
8.2.1. Модальность данных	210
8.2.2. Сценарии использования мультимодального RAG	210
8.2.3. Мультимодальные пайплайны RAG	211
8.2.4. Ключевые проблемы и лучшие практики	218
8.3. RAG с графом знаний	219
8.3.1. Графы знаний	219
8.3.2. Сценарии применения графового RAG	221
8.3.3. Подходы графового RAG	222
8.3.4. Пайплайны в графовом RAG	225
8.3.5. Проблемы и лучшие практики	229
8.4. Агентный RAG	230
8.4.1. LLM-агенты	230
8.4.2. Возможности агентного RAG	233
8.4.3. Пайплайны агентного RAG	234
8.4.4. Проблемы и лучшие практики	237
8.5. Другие варианты RAG	238
8.5.1. Корректирующий RAG	238
8.5.2. Спекулятивный RAG	240
8.5.3. Саморефлексивный RAG (Self-RAG)	240
8.5.4. RAPTOR	241
Резюме	243
Глава 9. Общая схема разработки RAG и направления для дальнейшего изучения	246
9.1. Общая схема разработки RAG	247
9.1.1. Этап инициации: определение и очерчивание границ RAG-системы	249
9.2. Этап проектирования: структурирование стека RAGOps	254
9.2.1. Проектирование пайплайна индексации	254
9.2.2. Проектирование пайплайна генерации	258
9.2.3. Дополнительные аспекты проектирования	262
9.2.4. Этап разработки: создание модульных пайплайнов RAG	263
9.2.5. Этап оценки: валидация и оптимизация RAG-системы	266

9.2.6. Этап развертывания: запуск и масштабирование RAG-системы	269
9.2.7. Этап поддержки: обеспечение надежности и адаптивности	271
9.3. Что изучать дальше	272
9.3.1. Дообучение в рамках RAG	272
9.3.2. Длинные контекстные окна в LLM	273
9.3.3. Управляемые решения	274
9.3.4. Трудные запросы	275
<i>Резюме</i>	275
Указатель	278

Предисловие

Каким образом машины понимают человеческие намерения? Этот вопрос всегда вызывал у меня глубокий интерес. Хотя мое путешествие в мир искусственного интеллекта (ИИ) и машинного обучения началось в 2007 году, по-настоящему я заинтересовался обработкой естественного языка (NLP, Natural Language Processing) в начале 2016-го, когда работал над созданием виртуального аналитика данных. Когда в 2018 году исследователи из Google представили языковую модель BERT, я убедился, что NLP находится на пороге настоящей революции.

В 2022-м, после релиза модели text-davinci-002 из серии GPT-3 от OpenAI, я решил присоединиться к Yarnit — платформе для контент-маркетинга на основе генеративного ИИ, чтобы разработать ее технологическую основу. Нашей целью было создать платформу, где корпоративные маркетинговые команды могли бы быстро, масштабно, с высокой точностью и минимальными затратами генерировать маркетинговые материалы: посты в соцсетях, блоги, электронные письма и многое другое. Вскоре стало очевидно, что без учета знаний о бренде и доступа к закрытым данным ни одна генеративная модель не справится с этой задачей сколь-нибудь эффективно. Это осознание подтолкнуло меня к изучению RAG (Retrieval-Augmented Generation, генерации с дополнением данными).

Большие языковые модели (LLM, Large Language Models) часто не оправдывают ожиданий пользователей. Они невероятно эффективны в хранении и генерации знаний, но при этом склонны к галлюцинациям — уверенно высказанным ошибочным выводам. Именно здесь RAG совершает прорыв, позволяя LLM получать релевантную, фактически точную и свежую информацию перед генерацией ответа. Прелесть RAG заключается в простоте концепции в сочетании с тонкостями реализации. Способность RAG радикально преодолевать ключевые ограничения LLM — вот что удерживает внимание как исследователей, так и практиков.

Когда я начал изучать RAG, эта область была еще мало исследована. Официальных материалов для изучения RAG практически не существовало, сведения были разбросаны по блогам, публикациям в соцсетях, научным статьям и дискуссионным форумам. Я активно делился собственными находками на социальных платформах и в блогах. Со временем у меня созрела идея объединить все эти знания в единую книгу.

Поставив перед собой цель создать простой и практичный ресурс для технических специалистов, разрабатывающих приложения на основе LLM, я приступил к работе над книгой в середине 2024 года. Со временем она превратилась в фундаментальное руководство по RAG, охватывающее тему как вширь, так и вглубь, с акцентом на практической реализации через понятные объяснения и лаконичный код на Python.

Я убежден, что знание RAG обязательно для каждого, кто работает с ИИ-приложениями, а его освоение требует прочной концептуальной базы. Эта книга призвана ее обеспечить. Работа над книгой дала мне невероятно ценный опыт, я многому научился, пока писал ее. Надеюсь, чтение моей книги окажется для вас столь же познавательным и увлекательным.

Благодарности

Чтобы написать книгу, требуются бесчисленные часы исследований и кропотливой работы, особенно когда речь идет о такой быстроразвивающейся области, как RAG, где новые исследования появляются почти каждую неделю. Эта книга не увидела бы свет без безграничной любви и поддержки моей жены Паллави (Pallavi). Ее поддержка и терпение помогли мне пройти этот путь, за что я ей бесконечно благодарен.

Я глубоко признателен своим сооснователям Джьотирмою (Jyotirmoy) и Акашу (Akash), а также всей команде Yarnit, которые внесли значительный вклад в мое понимание RAG. Практический опыт создания реальных ИИ-приложений, несомненно, обогатил эту книгу, сделав ее более ценной для читателей.

Также я хотел бы выразить искреннюю благодарность коллегам и наставникам: Ашишу Риши (Ashish Rishi), Сатьякаму Моханти (Satyakam Mohanty), Прадипте Мишре (Pradeepta Mishra), Мегхе Джон (Megha John), Сандипу Ачарье (Sandeep Acharya), Акшиту Шарме (Akshit Sharma), Вишалу Синхе (Vishal Sinha) и многим другим — за содержательные беседы и наставничество на протяжении многих лет. Их взгляды сформировали мою философию и подход к data science и ИИ.

Отдельная благодарность выдающейся команде Manning Publications во главе с Энди Валдроном (Andy Waldron) за предоставленную возможность.

Я глубоко признателен Яну Хафу (Ian Hough) за его бесценные советы в процессе написания книги. Особую признательность выражаю техническому редактору Артуро Гейгелю (Arturo Geigel) за тщательную проверку написанного и содержательные замечания, которые значительно улучшили книгу. Большое спасибо Азре Дедич (Azra Dedic) за существенное улучшение графики в книге, а также Робину Кэмпбеллу (Robin Campbell) и Аире Дуич (Aira Ducic) за их выдающуюся работу по продвижению и маркетингу книги. Также приношу благодарность сотрудникам производственного отдела за усердный труд по подготовке книги к публикации.

Я глубоко признателен сообществу исследователей ИИ, чье неустанное стремление к знаниям и инновациям продолжает расширять границы возможного. Во многом эта книга отражает коллективные знания исследователей, участников open-source проектов и практиков, щедро делившихся своими идеями в статьях, блогах и на форумах.

Я искренне признателен всем рецензентам: Абхишеку Гупте (Abhishek Gupta), Алехандро Куэвасу Риверо (Alejandro Cuevas Rivero), Алексу Маклинтоку (Alex McLintock), Алирезе Агамохаммади (Alireza Aghamohammadi), Амиту Дикситу (Amit Dixit), Анидите Нат (Anindita Nath), Аниндьядипу Санниграхи (Anindyadeep Sannigrahi), Арьяну Джадону (Aryan Jadon), Ашишу Саркару (Ashish Sarkar), Аушиму Нагаркатти (Aushim Nagarkatti), Авинашу Тивари (Avinash Tiwari), Баблу Кумару (Babloo Kumar), Баладжи Дхамодхарану (Balaji Dhamodharan), Балакришнану Баласубраманиану (Balakrishnan Balasubramanian), Берту Голлнику (Bert Gollnick), Бхаргобу Деке (Bhargob Deka), Брайану Дейли (Brian Daley), Чарану Акири (Charan Akiri), Кристоферу Дж. Фраю (Christopher G. Fry), Харчарану С. Каббаю (Harcharan S. Kabbay), Харшвардану С. Фартале (Harshwardhan S. Fartale), Игорю Свиленкову Божичу (Igor Svilenkov Božić), Ивану Морено (Iván Moreno), Лалиту Чоурей (Lalit Chourey), Луису Луангкесорну (Louis Luangkesorn), Луи-Франсуа Бушарду (Louis-François Bouchard), Манасу Талукдару (Manas Talukdar), Марсио Ф. Ногейре (Márcio F. Nogueira), Марин Серре (Marine Serré), Наре Сантошу Редди Вутукури (Naga Santhosh Reddy Vootukuri), Нилеш Патерия (Neelesh Pateriya), Питеру Котронео (Peter Cotroneo), Питеру Моргану (Peter Morgan), Риче Талдар (Richa Talдар), Ридхибен Суниткумар Шах (Riddhiben Sunitkumar Shah), Роберту Винсу (Robert Vince), Самиту Соनावане (Sameet Sonawane), Сашанку Даре (Sashank Dara), Стивену Вольфу (Stephen Wolff), Субхашу Кумару Периясами (Subhash Kumar Periasamy), Винешу Гудле (Vinesh Gudla) и Янки Луо (Yanqi Luo) — за их ценные идеи и предложения, которые помогли улучшить качество этой книги.

В заключение я хочу сердечно поблагодарить всех, кто был частью этого пути. Ваша поддержка, мудрость и щедрость помогли воплотить в жизнь не только эту книгу, но и мою мечту стать автором.

О книге

RAG трансформирует сферу прикладного генеративного искусственного интеллекта. Термин RAG впервые был предложен Льюисом (Lewis) и его коллегами в основополагающей работе Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks (<https://arxiv.org/abs/2005.11401>). Технология RAG быстро стала краеугольным камнем современного ИИ, повышая надежность и достоверность результатов больших языковых моделей (LLM).

«Простое руководство по дополненной генерации» — это фундаментальное пособие для всех, кто хочет изучить RAG. Книга представляет доступное, но всестороннее введение в концепцию, а также практические рекомендации, которые помогут эффективно использовать RAG в своих целях.

Для кого эта книга?

Книга предназначена для технических специалистов, которые хотят познакомиться с концепцией RAG и создавать приложения на основе LLM. Она будет полезна как новичкам, так и опытным профессионалам. Если вы специалист по данным, инженер данных, инженер машинного обучения, разработчик программного обеспечения, техлид или студент, интересующийся разработкой приложений на основе генеративного ИИ, вы найдете эту книгу полезной. После прочтения вы:

- разберетесь в основах, компонентах и практическом применении RAG;
- узнаете, как создавать непараметрические базы знаний и как с ними работать;

- сможете построить RAG-систему с глубоким погружением в пайплайн индексирования и генерации;
- получите полное представление об оценке RAG-систем и о модульных стратегиях оценки;
- познакомитесь с современными стратегиями и постоянно развивающимся ландшафтом RAG;
- приобретете знания об инструментах, технологиях и фреймворках для создания и развертывания систем RAG корпоративного уровня;
- узнаете о передовых вариантах RAG, таких как мультимодальный и агентный RAG;
- получите представление о текущих ограничениях RAG и познакомитесь с популярными новейшими техниками для их дальнейшего изучения.

Эта книга является фундаментальным руководством и не требует от читателя глубокого понимания концепций, однако будет полезно предварительно познакомиться с миром машинного обучения, генеративного ИИ и LLM. Вы глубже поймете суть LLM, прочитав первую главу.

Книга содержит примеры кода на Python с использованием фреймворка LangChain. Важно отметить, что эти фрагменты кода служат лишь дополнительной иллюстрацией концепций и предназначены для читателей, которые хотят получить практический опыт. Для работы с этим кодом достаточно базового понимания Python и API.

Генеративный ИИ остается областью прорывных, непрерывно развивающихся технологий. Эта книга поможет вам повысить квалификацию и открыть новые возможности в текущей и будущей профессиональной деятельности.

Как организована эта книга: краткий обзор

Книга состоит из девяти глав, распределенных по четырем частям.

Часть 1 посвящена основополагающим понятиям RAG.

- Глава 1 начинается с определения RAG и объяснения его необходимости и значимости в той области ИИ, где применяются LLM. Здесь вы также найдете несколько практических примеров применения систем, использующих технологию RAG.
- Глава 2 посвящена разбору основных компонентов RAG-системы. В ней представлены два ключевых пайплайна: пайплайн индексации и пайплайн генерации. Кроме того, в главе затрагиваются концепции оценки RAG.

Часть 2 посвящена созданию базовой RAG-системы с основными пайплайнами и их оценке.

- В главе 3 мы рассмотрим сквозной пайплайн индексации для создания базы знаний RAG-системы. Вы на примерах познакомитесь

с загрузкой данных, разбиением на чанки, эмбедингами (преобразованием текста в векторное представление) и векторным хранилищем.

- Глава 4 освещает пайплайн генерации, обеспечивающий доступ к базе знаний и взаимодействие с LLM для генерации контекстных и точных ответов в режиме реального времени. Мы поговорим о ретриверах, стратегиях извлечения релевантных данных и промпт-инжиниринге для RAG, а также получим общее представление о доступных LLM.
- В главе 5 подробно рассматриваются различные методы оценки RAG и проводится их анализ с точки зрения вопроса, ответа и контекста. Также мы обсудим создание эталонного набора данных (ground truth dataset) и его значимость. Из этой главы вы узнаете подробности о популярных фреймворках и бенчмарках для оценки RAG-систем.

Часть 3 предоставит руководство по улучшению RAG-пайплайна и расскажет о слоях, необходимых для построения RAG-системы, готовой для применения в продакшене.

- В главе 6 рассматриваются современные концепции RAG: наивный, продвинутый и модульный RAG. Мы поговорим о важных компонентах, а также о предпоисковых и постпоисковых стратегиях. В этой главе также представлены методы оптимизации для повышения производительности RAG-системы.
- Глава 7 знакомит с различными инструментами и технологиями, которые обеспечивают работу стека RAGOps (RAG Operations). Вы узнаете о критически важных слоях, без которых любая RAG-система не сможет функционировать, а также о слоях, необходимых для повышения производительности системы, и о дополнительных слоях, ориентированных на удобство использования, масштабируемость и эффективность.

В части 4 вы познакомитесь с популярными современными вариантами RAG и общей схемой разработки RAG.

- Глава 8 посвящена новейшим вариантам RAG, включая мультимодальный, графовый и агентный RAG.
- В завершающей, 9-й главе представлена общая схема разработки RAG, которая поможет вам спланировать создание RAG-системы.

Книга рассчитана на последовательное чтение, причем в последней главе дается обзор всех рассмотренных в книге концепций.

О коде

Все примеры кода в этой книге написаны на Python. Исполняемые фрагменты кода доступны в онлайн-версии книги на платформе liveBook по