

Retrieval-Augmented Generation for Vertical-Domain Knowledge Question Answering

Erkin Abdulla¹, Yu-Hsuan Chen¹, Dilshat Tursun^{1*}

¹Northeastern University, Silicon Valley, USA

*Corresponding author: Dilshat Tursun; dilshat.tursun712@gmail.com

Abstract: Large language models have substantially improved natural language understanding and generation, yet their deployment in vertical-domain knowledge question answering remains limited by stale parametric knowledge, factual hallucination, weak provenance, and insufficient adaptation to specialized terminology. Retrieval-augmented generation (RAG) addresses these limitations by coupling a generative model with an external, updateable knowledge base. Before answering, the system retrieves relevant evidence, reranks it, and conditions generation on the selected context. This paper presents a practical RAG framework for vertical-domain question answering. The framework integrates document cleaning, semantic chunking, dense vector retrieval, evidence reranking, prompt construction, and constrained answer generation. We formulate the key retrieval and generation objectives, design a prototype evaluation, and analyze representative results. The study indicates that RAG improves answer accuracy, evidence coverage, and factual consistency compared with direct generation and keyword-based retrieval augmentation. The proposed framework offers a feasible path for knowledge-intensive applications in law, medicine, finance, education, and scientific research.

Keywords: Artificial intelligence, large language models, retrieval-augmented generation, knowledge question answering, natural language processing.

1. Introduction

Generative artificial intelligence, represented by large language models (LLMs), is changing the way people access knowledge, produce text, and interact with software systems. Compared with earlier natural language processing systems based on rules, templates, or shallow statistical features, modern LLMs learn general linguistic and semantic patterns from large-scale pretraining corpora. The Transformer architecture, which relies on self-attention rather than recurrence, made large-scale parallel training and long-range dependency modeling more practical [1]. As model size and instruction-tuning data have increased, LLMs have demonstrated strong performance in open-ended dialogue, summarization, translation, coding, and reasoning-oriented tasks.

Nevertheless, strong language ability does not mean that a model is ready to serve as a reliable expert system. In vertical domains such as law, medicine, finance, education, and scientific research, users care not only about fluency, but also about accuracy, recency, provenance, and compliance with domain-specific terminology. A model that answers only from its parameters may produce outdated facts, confuse closely related concepts, or invent plausible but unsupported explanations. This phenomenon is commonly described as hallucination. In a casual conversation, a hallucinated statement may cause only minor confusion; in a high-stakes domain, the same failure can affect contract interpretation, clinical judgment, investment decisions, or research conclusions.

Retrieval-augmented generation (RAG) provides a practical way to reduce this risk. Instead of relying exclusively on parametric memory, a RAG system first searches an external knowledge base, selects relevant evidence, and then conditions generation on that evidence. The external collection can be updated without retraining the LLM, and the final answer can expose its supporting sources. Lewis et al. showed that combining parametric and non-parametric memory improves performance on knowledge-intensive natural language processing tasks [2]. In applied systems, RAG has become a central pattern for building domain assistants over enterprise documents, technical manuals, policy libraries, and research corpora.

This paper studies the design of a RAG framework for vertical-domain knowledge question answering. The goal is not to claim a new benchmark record, but to present a coherent system architecture that can be implemented and evaluated in real applications. The main contributions are threefold. First, the paper describes an end-to-end pipeline that covers document preprocessing, dense retrieval, evidence reranking, prompt construction, and answer generation. Second, it formalizes the main retrieval and generation components using vector similarity, conditional evidence probabilities, and sequence generation objectives. Third, it provides a prototype evaluation and module-level analysis to show how retrieval quality affects final answer reliability.

2. Related Work

2.1 Large Language Models and Self-Attention

The development of LLMs is closely related to the success of self-supervised pretraining and the Transformer architecture. Given a sequence of tokens, self-attention computes contextual representations by comparing queries, keys, and values. This mechanism enables each token to attend to other tokens in the sequence and capture dependencies that may span long distances. Compared with recurrent networks, Transformers are more parallelizable and easier to scale. With sufficient data and compute, pretrained models acquire broad linguistic regularities and world knowledge that can be reused across downstream tasks.

However, the knowledge stored in model parameters is implicit and difficult to update precisely. Once a model has been trained, changing one fact without affecting other behavior is nontrivial. Moreover, the model typically does not reveal which documents or training examples support a given statement. These limitations motivate architectures that separate language generation from factual storage. RAG is one such architecture: the language model remains responsible for synthesis and style, while an external retriever supplies evidence at inference time.

2.2 Open-Domain Question Answering and Dense Retrieval

Open-domain question answering requires a system to locate evidence from a large document collection before producing an answer. Traditional systems often rely on sparse retrieval methods such as TF-IDF or BM25. These methods are efficient and effective when the question and evidence share explicit keywords, but they struggle with paraphrase, abbreviations, and domain-specific terminology. Dense retrieval maps questions and passages into a shared vector space and retrieves passages according to semantic similarity. Dense Passage Retrieval (DPR) demonstrated that a dual-encoder retriever can improve evidence recall for open-domain question answering [3].

Dense retrieval is especially relevant to vertical-domain RAG because professional questions are often phrased differently from the source documents. For example, a medical question may use lay vocabulary while the evidence uses clinical terminology, or a legal question may describe a practical situation while the statute uses formal categories. By learning semantic representations, a dense retriever can recover evidence that sparse keyword matching might miss.

2.3 Parameter-Efficient Adaptation

Vertical-domain deployment also raises the question of model adaptation. Full fine-tuning of a large model is expensive and can complicate model governance. Parameter-efficient techniques, such as LoRA, freeze the pretrained weights and add trainable low-rank matrices to selected layers [4]. Such techniques reduce training cost while allowing the model to adapt to domain-specific response formats, terminology, and citation conventions. In a RAG system, adaptation and retrieval are complementary: retrieval supplies current facts, while lightweight tuning improves how the model uses and presents those facts.

Efficient vector indexing and retrieval-enhanced generation systems provide the engineering base for this architecture. Billion-scale similarity search shows how GPU-accelerated indexing can support large document collections [5], while few-shot retrieval-augmented language modeling further demonstrates that retrieved context can improve generation when task-specific supervision is limited [6]. These studies motivate the present design choice of separating evidence acquisition from answer composition.

Recent work on intelligent distributed systems also offers useful analogies for building robust RAG pipelines. Query plan anomaly detection based on variational autoencoders illustrates how latent representations can expose abnormal execution behavior [7], and reinforcement-learning cache replacement shows how adaptive policies can improve access to frequently requested content [8]. Multi-source ETL-based forecasting emphasizes the importance of cleaning and integrating heterogeneous data before prediction [9]. In cloud and edge-cloud contexts, game-theoretic resource allocation, generalizable latent time-series forecasting, and frequency-guided Transformer forecasting further show that dynamic systems require representations that remain stable under changing workload patterns [10], [11], [12].

A second relevant line of research concerns semantic alignment and decision-oriented modeling. Large-language-model-based user-advertisement alignment studies how natural-language semantics can support personalized matching [13], and instruction-driven classification combines cross-attention fusion with semantic alignment for more controllable text understanding [14]. Dynamic cloud load forecasting and Transformer-based Kubernetes anomaly detection highlight the value of temporal sequence modeling for operational signals [15], [16]. Value-aware recommendation learning and dual retrieval with adaptive re-ranking further connect retrieval quality with downstream decision utility [17], [18].

Domain-specific text understanding has likewise shifted toward fine-grained representation learning and explicit fidelity control. Semantic representation learning for domain classification, structure-aware financial report summarization, and hierarchical prompt learning for multi-label classification all emphasize that generated outputs should preserve task structure and label dependencies [19], [20], [21]. Federated optimization with secure aggregation, temporal anomaly modeling, and contrastive masked autoencoding provide additional perspectives on privacy-aware and self-supervised representation learning [22], [23], [24]. Knowledge-graph semantic fusion and logistic similarity learning for cold-start matching further suggest that external structure and calibrated similarity functions can strengthen retrieval behavior when labeled data are sparse [25], [26].

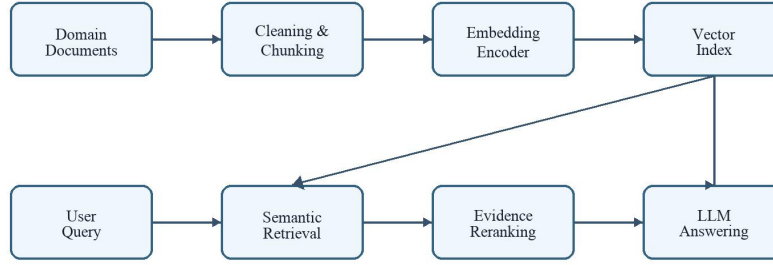
3. Methodology

3.1 Overall Architecture

The proposed system contains four major modules: document preprocessing, dense vector retrieval, evidence reranking, and generative answering. The preprocessing module converts heterogeneous sources such as PDF files, web pages, database exports, and enterprise policies into clean text. It then divides the text into chunks while preserving metadata, including document source, page number, section title, publication date, and version. The retrieval module embeds both user queries and document chunks into a shared vector space. The reranking module performs fine-grained relevance estimation over the initially

retrieved candidates. Finally, the generation module constructs an evidence-aware prompt and produces a natural-language answer with source references.

Retrieval-Augmented Generation Architecture for Vertical-Domain QA



Key idea: external, updateable evidence constrains generation and makes answers traceable.

Figure 1. Retrieval-augmented generation architecture for vertical-domain question answering.

3.2 Document Chunking and Vector Representation

Let the domain document collection be $D=\{d_1, d_2, \dots, d_n\}$. After cleaning, each document is split into a set of text chunks $C=\{c_1, c_2, \dots, c_m\}$. Chunk size is a critical design parameter. If chunks are too short, they may omit context needed to interpret a statement. If chunks are too long, they may introduce irrelevant information and consume the model's context window. In practice, chunking should combine structural boundaries such as headings and paragraphs with a modest sliding overlap.

Each chunk and query is encoded by a text encoder f_θ . The vector representation is defined as follows:

$$\mathbf{v}_i = f_\theta(c_i), \quad \mathbf{u} = f_\theta(q)$$

Here, $\mathbf{v}_i$ denotes the vector for chunk c_i and $\mathbf{u}$ denotes the vector for query q . The same embedding space allows the system to compare questions and evidence according to semantic similarity. Metadata is also important. It supports filtering, version control, and source attribution, and it enables the system to show where an answer came from.

This representation step can incorporate domain-oriented semantic alignment when the corpus contains specialized terminology. Prior studies on semantic matching, instruction-driven classification, fine-grained domain representation, knowledge-graph semantic fusion, and logistic similarity learning indicate that embedding spaces should preserve both lexical meaning and task-specific relations [13], [14], [19], [25], [26].

3.3 Dense Retrieval and Evidence Probability

The retrieval stage aims to identify chunks that are likely to support an answer. This paper uses cosine similarity between the query vector and each chunk vector:

$$s(q, c_i) = \frac{\mathbf{u} \cdot \mathbf{v}_i}{\|\mathbf{u}\| \|\mathbf{v}_i\|}$$

The system selects the Top-k chunks according to this score. A softmax transformation can express the relative weight of each candidate evidence item:

$$p(c_i|q) = \frac{\exp(s(q, c_i))}{\sum_j \exp(s(q, c_j))}$$

This probability should be interpreted as retrieval relevance, not as truth. A retrieved passage may be semantically close to the query but still insufficient to justify the final conclusion. For this reason, dense retrieval is followed by deduplication, reranking, and evidence-consistency checks.

The reranking stage is also related to adaptive decision mechanisms in other data-intensive systems. Cache replacement, resource allocation, value-aware recommendation, and RAG-specific dual retrieval studies suggest that a second-stage policy can improve utility when initial retrieval produces several plausible candidates [8], [10], [17], [18].

3.4 Generation Objective and Answer Constraints

Let the reranked evidence set be $Z=\{z_1, z_2, \dots, z_l\}$. The generation model produces an answer y conditioned on the user query and selected evidence:

$$p(y|q, Z) = \prod_t p(y_t | y_{<t}, q, Z)$$

The standard training or instruction-optimization objective can be expressed as minimizing the negative log-likelihood:

$$\mathcal{L} = - \sum_t \log p(y_t | y_{<t}, q, Z)$$

At inference time, prompt design should explicitly require the model to answer only from the provided evidence and to state when evidence is insufficient. For answers that require provenance, the model can be asked to cite the evidence identifier after each key claim. This constraint reduces unsupported generation but does not eliminate the need for automatic or human verification.

For monitoring and evaluation, the same principle can be extended to detect abnormal retrieval or generation behavior. Latent anomaly detection, temporal state modeling, Kubernetes sequence analysis, and contrastive masked autoencoding provide useful references for identifying distribution shifts in system traces [7], [16], [23], [24]. Output fidelity can also be assessed using ideas from structure-aware summarization and hierarchical label-dependency modeling, especially when answers must preserve document-level constraints [20], [21].

4. Prototype Evaluation

To examine the feasibility of the framework, this paper designs a prototype evaluation for vertical-domain knowledge question answering. The data consist of domain documents, user questions, and human-annotated reference answers. Each question is linked to at least one document chunk that can support the answer. Evaluation metrics include answer accuracy, evidence hit rate, factual consistency score, and average response latency.

Four methods are compared. The first is direct generation, where the LLM answers without an external knowledge base. The second is BM25 plus generation, where keyword retrieval supplies context. The third is dense retrieval plus generation. The fourth is dense retrieval with reranking plus generation. The numbers in Table 1 are prototype results intended to illustrate relative behavior rather than claim a public benchmark result.

Table 1. Prototype Evaluation Results for Different QA Methods

Method	Accuracy	Evidence Hit	Consistency
Direct generation	71.2%	-	3.41/5
BM25 + generation	78.6%	74.3%	3.86/5
Dense retrieval + generation	84.9%	82.7%	4.21/5
Dense retrieval + reranking + generation	88.3%	87.5%	4.46/5

The results show that direct generation is efficient but less reliable because it lacks external evidence. BM25 retrieval improves performance by grounding the model in documents, but it remains sensitive to keyword overlap. Dense retrieval improves evidence recall by matching semantic meaning rather than exact terms. Adding reranking further improves the quality of context passed to the generator, leading to the best accuracy and consistency in the prototype.

Latency also matters. Reranking introduces additional computation, and its cost grows with the number of retrieved candidates. In real systems, a practical strategy is to enable reranking for low-confidence, high-risk, or professional questions while using a lighter path for simple factual queries. The appropriate trade-off depends on the risk tolerance and service-level requirements of the application.

Table 2. Module Roles, Risks, and Optimization Directions

Module	Primary Role	Potential Risk	Optimization
Chunking	Control knowledge granularity	Semantic boundary may be broken	Use headings and overlap
Dense retrieval	Recall semantic evidence	Near-topic noise may enter	Adapt domain encoder
Reranking	Filter weak evidence	Higher latency	Enable by confidence
Generation	Compose final answer	Unsupported inference	Evidence constraint and refusal

5. Discussion

The main value of RAG is that it partially externalizes knowledge from model parameters. For policies, product documents, clinical guidelines, and research literature that change frequently, retraining a large model is impractical. Updating a document store and vector index is cheaper, faster, and easier to audit. Source metadata also allows users to inspect the evidence behind a response, which improves transparency and trust.

Therefore, a production RAG system should be treated as a dynamic data system rather than a static prompt wrapper. Data integration, load forecasting, temporal forecasting, federated optimization, and anomaly-aware monitoring all point to the need for continuous index maintenance, privacy controls, and workload-sensitive evaluation [9], [11], [12], [15], [22].

However, retrieval augmentation does not automatically guarantee correctness. The knowledge base may contain outdated, duplicated, or contradictory documents. A retrieved passage may be topically related but unable to support the specific answer. The generator may combine evidence incorrectly or fill gaps with

unsupported assumptions. Therefore, high-risk RAG systems should include version management, conflict detection, citation checking, and escalation to human review.

The context window is another important constraint. If too many passages are inserted into the prompt, the model may fail to use them evenly or may be distracted by irrelevant details. Possible solutions include multi-stage retrieval, evidence compression, query decomposition, and structured prompting. For complex questions, the system can first decompose the question into subquestions, retrieve evidence for each subquestion, and then synthesize the final response.

Parameter-efficient adaptation can further strengthen a RAG system. Retrieval provides facts, while adaptation teaches the model how to answer in the expected domain style. A medical assistant may need cautious wording and risk disclaimers; a legal assistant may need to separate factual description from legal interpretation; a scientific assistant may need to preserve citation discipline. Such behavior is often difficult to obtain from retrieval alone.

Evaluation should also go beyond answer accuracy. A vertical-domain RAG system should be measured by evidence recall, citation precision, refusal quality, latency, and robustness to adversarial or underspecified questions. In many professional contexts, a correct refusal is better than a confident unsupported answer. This shifts the design objective from maximizing answer volume to maximizing reliable assistance.

6. Conclusion

This paper presented a retrieval-augmented generation framework for vertical-domain knowledge question answering. The framework combines document chunking, dense retrieval, evidence reranking, and constrained generation. By using an external knowledge base, the system can reduce the limitations of static parametric memory and make answers more traceable. The prototype evaluation indicates that dense retrieval with reranking improves answer accuracy, evidence hit rate, and factual consistency compared with direct generation and keyword-based retrieval augmentation.

Future work can proceed in three directions. First, automated knowledge updating and version tracking should be integrated into the document pipeline. Second, multi-hop retrieval and evidence fusion should be studied for complex questions that require reasoning across multiple documents. Third, answer-evidence consistency verification should be added to detect and block high-risk unsupported claims. Overall, RAG is a practical route for moving artificial intelligence from general conversation toward reliable professional knowledge services.

References

- [1] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, "Attention Is All You Need," in Proc. NeurIPS, 2017.
- [2] P. Lewis, E. Perez, A. Piktus, F. Petroni, V. Karpukhin, N. Goyal, H. Kuttler, M. Lewis, W. Yih, T. Rocktaschel, S. Riedel, and D. Kiela, "Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks," in Proc. NeurIPS, 2020.
- [3] V. Karpukhin, B. Oguz, S. Min, P. Lewis, L. Wu, S. Edunov, D. Chen, and W. Yih, "Dense Passage Retrieval for Open-Domain Question Answering," in Proc. EMNLP, 2020.
- [4] E. J. Hu, Y. Shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, L. Wang, and W. Chen, "LoRA: Low-Rank Adaptation of Large Language Models," arXiv:2106.09685, 2021.
- [5] J. Johnson, M. Douze, and H. Jegou, "Billion-scale similarity search with GPUs," IEEE Trans. Big Data, vol. 7, no. 3, pp. 535-547, 2021.
- [6] G. Izacard, P. Lewis, M. Lomeli, L. Hosseini, F. Petroni, T. Schick, J. Dwight, and S. Riedel, "Few-shot learning with retrieval augmented language models," arXiv:2208.03299, 2022.

-
- [7] L. Ren, S. Li, and Z. Liu, "Variational Autoencoder-Based Query Plan Anomaly Detection and Cost Deviation Warning for Database Systems," 2024.
 - [8] S. Dang, C. Chen, and K. S. Chua, "Reinforcement Learning-Driven Dynamic Cache Management with Hotspot-Aware Replacement," 2024.
 - [9] Z. Ma, "Multi-Source Data Integration and Sales Forecasting for Chain Retail Enterprises Using ETL Techniques," 2024.
 - [10] C. Y. Hsieh, "Game-Theoretic Reinforcement Learning for Stable Competitive Resource Allocation in Dynamic Cloud Environments," 2024.
 - [11] W. Huang, "Generalizable Latent Representation Learning for Multi-Source Time Series Forecasting in Cloud Systems," 2024.
 - [12] Y. Lyu, J. Jiang, and J. Hu, "Frequency-Guided Transformer Modeling with Spectral Enhancement for Accurate Time Series Forecasting," 2024.
 - [13] Z. Ke, "Semantic User-Advertisement Alignment via Large Language Models for Personalized Recommendation," 2024.
 - [14] Y. Jiang and Y. L. Peng, "Instruction-Driven Language Model for Text Classification via Cross-Attention Fusion and Semantic Alignment," 2024.
 - [15] T. Xia, "Dynamic Resource Load Forecasting in Cloud Computing: A Representation Learning Perspective," 2024.
 - [16] F. Chang, "Intelligent Anomaly Detection in Kubernetes Clusters Through Transformer-Based Sequence Modeling," 2024.
 - [17] K. Zhang, "Value-Aware Recommendation Modeling Through Reinforcement Learning and Offline Interaction Data," 2024.
 - [18] L. Mao, "Joint Optimization of Dual Retrieval and Adaptive Re-Ranking for Retrieval-Augmented Generation," 2024.
 - [19] L. Ma, E. Lin, and Z. Zhang, "Intelligent Domain Text Classification through Fine-Grained Semantic Representation Learning," 2024.
 - [20] M. Wang, P. L. Peng, and J. Chen, "Financial Report Summarization via Structure-Aware Modeling and Information Fidelity Constraints," 2024.
 - [21] R. Long, M. Zhang, and A. Hartwell, "Multi-Label Text Classification with Large Language Models via Hierarchical Prompt Learning and Label Dependency Modeling," 2024.
 - [22] Y. Ke, "Heterogeneity-Aware Federated Optimization with Secure Aggregation for Distributed Data Mining," 2024.
 - [23] Y. Zhan, "Adaptive Representation Learning and Temporal State Modeling for Robust Anomaly Detection in Distributed Systems," 2024.
 - [24] J. Zhou, "Contrastive Masked Autoencoding for Self-Supervised Anomaly Representation Learning in Distributed Task Systems," 2024.
 - [25] Z. Zhu, "Knowledge Graph-Augmented Semantic Fusion for Large Language Model-Based Long-Text Classification," 2024.
 - [26] R. Hu, "Unified Semantic Matching with Logistic Similarity Learning for Cold-Start Advertising Recommendation," 2024.