

Unlocking Multi-View Insights in Knowledge-Dense Retrieval-Augmented Generation

Anonymous ACL submission

Abstract

While Retrieval-Augmented Generation (RAG) plays a crucial role in the application of Large Language Models (LLMs), existing retrieval methods in knowledge-dense domains like law and medicine still suffer from the insufficient utilization of multi-perspective views embedded within domain-specific corpora, which are essential for improving interpretability and reliability. Previous research on multi-view retrieval often focused solely on different semantic forms of queries, neglecting the expression of specific domain knowledge perspectives. This paper introduces a novel multi-view RAG framework, *MVRAG*, tailored for knowledge-dense domains, which leverages machine learning techniques for professional perspectives extraction and intention-aware query rewriting from multiple domain viewpoints to enhance retrieval precision, thereby improving the effectiveness of the final inference. Experiments conducted on both retrieval and generation tasks demonstrate substantial improvements in generation quality while maintaining retrieval performance in complex, knowledge-dense scenarios.

1 Introduction

In the ever-evolving domain of natural language processing, retrieval-augmented generation (RAG) stands as a cornerstone technology that synergistically combines large language models (LLMs) with a vast corpus of external knowledge (Gao et al., 2023). This integration endows RAG systems with a remarkable capacity for nuanced language understanding and sophisticated information retrieval. Such capabilities are especially transformative in knowledge-dense domains like *law*, *medicine* and *biology*, where RAG not only augments contextual comprehension but also improves the interpretability and reliability of the models. Despite these advantages, current RAG implementations face substantial challenges in adequately serving

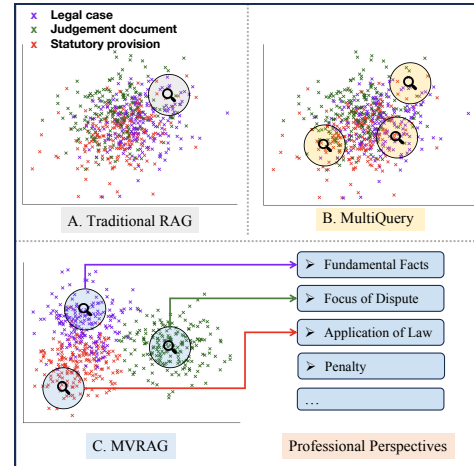


Figure 1: A t-SNE visualization of retrieval results from different methods using a legal database with manually added category labels, displayed in different colors in the plot. The magnifying glass represents the query vector, while the circles represent the retrieval results.

the complicated and multifaceted information requisites inherent in these specialized fields.

Predominant among these challenges is the insufficient utilization of multi-view information embedded within domain-specific corpora. Unlike general-purpose corpora such as Wikipedia, domain-specific corpora contain self-organizing structural information, which is reflected in the distribution patterns of vectors within the vector space they form. These patterns represent the embedded *professional perspectives* of the domain, guiding deeper and more thorough utilization of the corpus. These perspectives often play a more significant role than superficial textual similarity. In the legal domain, for example, the *focus of dispute* between cases is more crucial than mere textual overlap; in medical scenarios, the relevance of *medical history* or *symptoms* is often more diagnostically valuable than literal similarity, which can be obscured by irrelevant information.

Traditional retrieval methods (Gao et al., 2023),

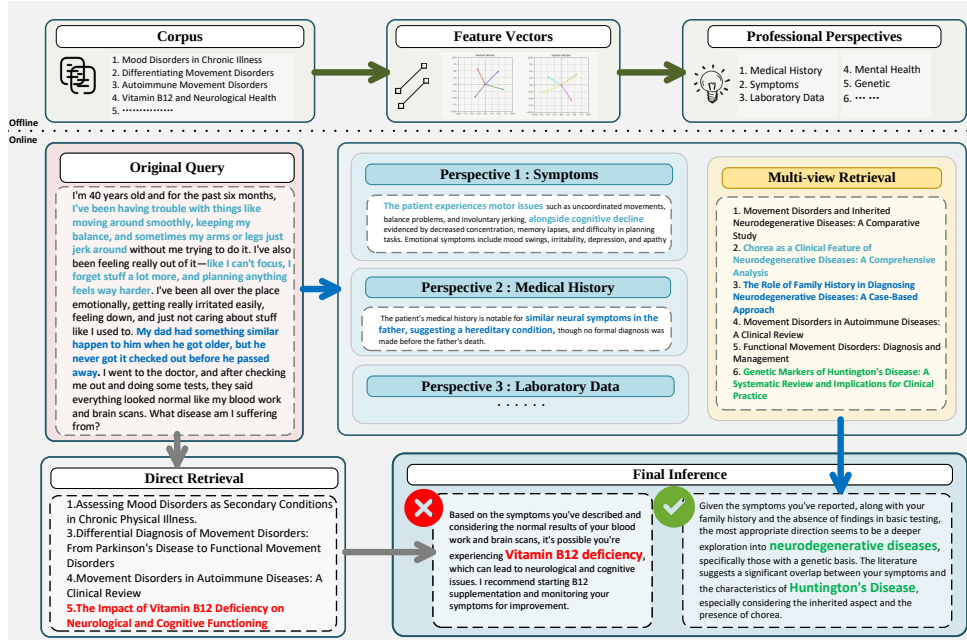


Figure 2: A case study showcasing the effectiveness of a multi-view retrieval framework in accurately diagnosing Huntington’s disease. The model corrects an initial misdiagnosis of Vitamin B12 deficiency by refining the search criteria to focus on neurodegenerative symptoms and family medical history, thus demonstrating the importance of multi-view search strategies in medical diagnostics. The professional perspectives used in the framework were determined during the offline part, ensuring their domain-specific relevance. Detailed case study and additional examples from other domains are provided in Appendix D and Appendix E.

as shown in Figure 1A, rely on full-text similarity, retrieving only a limited subset of vectors around the query vector. This often misses essential professional nuances, leading to imprecise or incomplete results (Wang et al., 2023). Current multi-view retrieval research¹ addresses this by rephrasing queries to adjust their vector space position (Figure 1B) (Ma et al., 2023), but this introduces randomness and fails to meet the complexity of professional knowledge domains. Figure 2 highlights how the lack of multi-view insights can lead to erroneous outcomes, emphasizing the need for a more nuanced and context-aware retrieval approach.

To address these challenges and overcome the limitations of existing methods, we propose a novel multi-view retrieval framework, MVRAG. This system is built upon three key stages: 1) professional perspectives extraction, 2) intention recognition and query rewriting, and 3) retrieval augmentation. For a specialized corpus, we first use machine learning techniques like PCA to extract domain-specific perspectives. The query’s intention is identified by measuring its similarity to these perspectives, followed by rewriting the query for each perspective

using an LLM. Results are retrieved and re-ranked based on perspective importance and retrieval similarity, then fed into the LLM generator. This process, illustrated in Figure 1C, distributes retrieval outcomes across multiple professional perspectives, enabling comprehensive corpus utilization.

We tested MVRAG on legal and medical qualification exams, where it significantly outperformed traditional RAG, especially on challenging multiple-choice and subjective questions. The contributions of this article mainly include the following three aspects:

- We present MVRAG, a framework that integrates professional perspectives extraction, intention-aware multi-view query rewriting, and retrieval re-ranking to enable multi-perspective analysis in knowledge-intensive domains. This system is designed for easy integration and introduces a new paradigm for applying RAG in specialized fields.
- We propose an innovative approach that incorporates traditional machine learning methods into the RAG framework for feature extraction from domain-specific corpora. By combining these methods with query rewriting tech-

¹https://python.langchain.com/v0.2/docs/how_to/MultiQueryRetriever

niques, this approach enhances both retrieval accuracy and depth.

- We conducted experiments on both retrieval and generation tasks, highlighting the importance of multi-perspective strategies in knowledge-dense RAG tasks, validating their effectiveness and paving the way for future research.

2 Related Work

2.1 Retrieval-Augmented Generation (RAG)

Retrieval-augmented generation, commonly known as RAG (Izacard et al., 2022; Huo et al., 2023; Guu et al., 2020; Lewis et al., 2020) has become a prevalent technique to enhance LLMs. This approach integrates LLMs with retrieval systems, enabling them to access domain-specific knowledge and base their responses on factual information (Khat-tab et al., 2021). Additionally, RAG provides a layer of transparency and interpretability by allowing these systems to cite their sources (Shuster et al., 2021).

Recent studies have demonstrated enhancements in the quality of output results across various scenarios through techniques such as appending documents retrieved by RAG to the inputs of LLMs, and training unified embedded models. These methods have shown effective improvements in the performance of LLM outputs by leveraging RAG’s ability to access and incorporate relevant information from external sources (Ram et al., 2023; Es et al., 2023).

2.2 Domain-Specific Large Language Models

Due to the general nature of LLMs, their expertise in specific domains such as *law*, *finance*, and *health-care* is often limited. Recent research focuses on enhancing domain-specific expertise through Knowledge Enhancement techniques, introducing specific domain knowledge, and employing innovative training methods to address the issue of hallucinations in models. This approach has become a key strategy for improving the professionalism of vertical domain large models (Xi et al., 2023; Yao et al., 2023; Zhao et al., 2023; Zhu et al., 2023).

In the legal field, a number of well-known large language models have been born through methods such as secondary training, instruction fine-tuning, and RAG, such as wisdomInterrogatory², DISC-

LawLLM³, PKUlaw⁴, and ChatLaw (Cui et al., 2023). These models have been specifically designed to address the unique requirements of the legal domain, providing specialized knowledge and expertise to enhance the quality and relevance of legal information retrieval and generation.

As for the medical domain, the development of domain-specific LLMs has been a key focus in recent research. Models like Huatuo⁵, Zhongjing⁶, and Doctor-Dignity⁷ have been specifically designed to cater to the complex and nuanced requirements of the medical field, providing enhanced capabilities for medical information retrieval, diagnosis, and treatment planning.

2.3 Query Rewriting

In the process of RAG, challenges often arise from users’ original queries being imprecisely phrased or lacking sufficient semantic information. Direct searches with such queries may lead Large Language Models to provide inaccurate or unanswerable responses. Thus, aligning the semantic space of user queries with that of document semantics is crucial. Query Rewriting techniques address this issue by refining the expression of queries, making them more precise and enriched, effectively bridging the gap and enhancing the accuracy and relevance of retrievals.

Gao et al. (2022) introduce an innovative approach predicated on the utilization of conjectural document embeddings, designed to maximize the congruence between query vectors and the corresponding authentic document representations, thereby augmenting the precision of query retrieval mechanisms. Concurrently, Wang et al. (2023) advocate for the employment of LLM as auxiliary tools in query formulation processes. Through the methodologies of query rewriting and the generation of pseudo-documents, among other modifications, their strategies have been empirically validated to deliver notable enhancements in retrieval outcomes (Ma et al., 2023).

Different from previous research, our work concentrates on the multi-perspective professional information within specific domains, which is a critical aspect often overlooked before. Instead of

²<https://github.com/zhihaiLLM/wisdomInterrogatory>

³<https://github.com/FudanDISC/DISC-LawLLM>

⁴<https://www.pkulaw.net/>

⁵<https://github.com/SCIR-HI/Huatuo-Llama-Med-Chinese>

⁶<https://github.com/SupritYoung/Zhongjing>

⁷<https://github.com/llSourceCell/Doctor-Dignity>

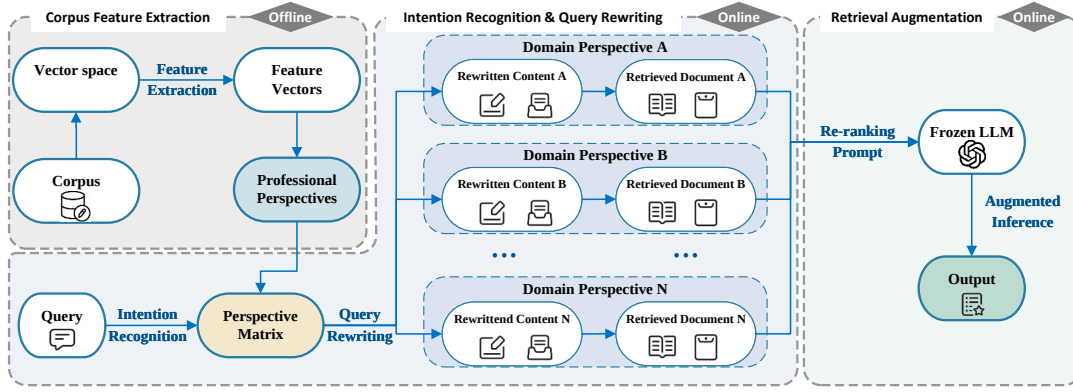


Figure 3: Framework of our Multi-View RAG System. This figure demonstrates the system’s core processes: Professional Perspectives Extraction, Intention Recognition and Query Rewriting, and Retrieval Augmentation, emphasizing the multi-view insights approach for intention-aware query rewriting

merely adjusting queries for semantic alignment, our novel multi-view retrieval framework captures the complex relationships and local nuances inherent to each domain, enhancing the retrieval process by focusing on the multi-perspective aspects of domain-specific knowledge.

3 Method

3.1 Overview

The Multi-View Retrieval-Augmented Generation (MVRAG) system comprises three main components: (1) **Professional Perspectives Extraction**, (2) **Intention Recognition & Query Rewriting**, and (3) **Retrieval Augmentation**.

Professional Perspectives Extraction (Offline): The domain-specific corpus is embedded into a vector space using a tokenizer model. Techniques such as Non-negative Matrix Factorization (NMF) and Principal Component Analysis (PCA) are applied to extract *professional perspectives*, representing key viewpoints in the corpus. This offline step provides reusable domain representations for later stages.

Intention Recognition & Query Rewriting (Online): A new query is vectorized and aligned with the extracted perspectives, producing a *Perspective Vector* that reflects its relevance to each perspective. Based on this alignment, the query is rewritten for each relevant perspective using a rewriter model. Rewritten queries are used for similarity-based document retrieval.

Retrieval Augmentation (Online): Retrieved documents are re-ranked by combining similarity scores with perspective importance weights. The top-ranked documents are integrated into a struc-

tured prompt, enabling the LLM to generate a contextually rich, multi-perspective response.

This pipeline integrates offline preprocessing with real-time query handling, ensuring efficient and contextually nuanced responses. Figure 3 illustrates the overall workflow, with detailed prompt templates provided in Appendix C.

3.2 Professional Perspectives Extraction

In detail, professional perspectives are topic terms that represent various aspects of a specific domain. These topic terms are then used to rewrite the query as is described in Section 3.3. In this section, we first capture the vector-level structural patterns in the corpus by PCA (**Feature Vector Extraction**), and then we employ NMF to identify candidate lists of (word-level) perspective-related topic terms (**Topic Modeling**). Finally, we translate the **feature vectors** into the **topic terms** (**Alignment of Perspectives**).

Feature Vector Extraction. Each document $d_i \in C$ in the corpus C is embedded as a vector $\mathbf{v}_i \in \mathbb{R}^n$ in a high-dimensional vector space generated using a domain-adapted embedding model that aligns with the LLM employed in downstream tasks. To uncover the underlying structure of the corpus and reduce dimensionality, Principal Component Analysis (PCA) ⁸ is applied:

$$\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n = \text{PCA}(C), \quad (1)$$

where $\mathbf{v}_i$ represents the i -th principal component, capturing key structural patterns of the corpus. The

⁸PCA was implemented using the scikit-learn library (Pedregosa et al., 2011), version 1.6.0. For more details, see <https://scikit-learn.org/stable/>.

Queries	Professional Perspectives						
	Symptoms	Medical History	Laboratory Data	Physical Examination	Emotional Factors	Medication Use	...
"I've been feeling extremely tired lately, and my skin is dry. Could this be related to my thyroid?"	0.49	0.31	0.25	0.17	0.05	0.23	...
"My father was diagnosed with Type 2 diabetes; should I get tested too?"	0.10	0.53	0.47	0.03	0.02	0.25	...
"I have a sharp pain on the right side of my abdomen. What could it be?"	0.77	0.19	0.23	0.69	0.05	0.02	...

Figure 4: Visualization of Perspective Vectors, using varying shades of color to represent different normalized scores. The graph displays various scenarios corresponding to different types of queries.

number of principal components n is determined by retaining 90% of the cumulative variance, ensuring a balance between computational efficiency and the preservation of key information.

Topic Modeling. The corpus is transformed into a document-term matrix $X \in \mathbb{R}^{|C| \times m}$ using Term Frequency-Inverse Document Frequency (TF-IDF), which represents term relevance across documents. Here, m denotes the total number of unique terms in the corpus, representing the dimensions of the term space. Non-negative Matrix Factorization (NMF) is then applied to X , yielding two non-negative matrices:

$$X \approx WH, \quad (2)$$

where $W \in \mathbb{R}^{|C| \times r}$ encodes document-topic relationships, and $H \in \mathbb{R}^{r \times m}$ encodes topic-term distributions. The number of topics r is chosen by maximizing coherence scores on a validation dataset, ensuring optimal interpretability. For each topic j , the top k terms from the corresponding row $\mathbf{h}_j$ of H are extracted:

$$S = \bigcup_{j=1}^r \text{Top}_k(\mathbf{h}_j), \quad (3)$$

where S is the set of candidate terms representing each topic. The number of terms k is fixed at 10 to provide a consistent representation of topics.

Notably, topic terms can be individual words or phrases, determined by the tokenization method in TF-IDF. In subsequent Chinese-context experiments, translation may cause notable variations in term lengths.

Alignment of Perspectives. To align extracted topics with the structural components, cosine similarity is computed between each principal component $\mathbf{v}_k$ and the topic term s_j in the obtained

set S from NMF. The topic terms are embedded using the same embedding model to ensure alignment with the dimensionality of $\mathbf{v}_i$. Each principal component $\mathbf{v}_i$ is assigned the topic terms s_j that maximizes the similarity:

$$p_i = \arg \max_{s_j} \frac{\mathbf{v}_i \cdot \mathbf{s}_j}{\|\mathbf{v}_i\| \|\mathbf{s}_j\|}, \quad (4)$$

where p_i represents the professional perspective corresponding to $\mathbf{v}_i$ and $\mathbf{s}_j$ represents the embedded vector of the topic term $s_j \in S$.

As a result of this process, we obtain a set of professional perspectives $P = \{p_1, p_2, \dots, p_n\}$, where each p_i encapsulates a distinct combination of structural and semantic patterns in the corpus.

3.3 Intention Recognition and Query Rewriting

After completing the structural analysis and feature extraction of the corpus, we obtain a set of professional perspectives $P = \{p_1, p_2, \dots, p_n\}$, where each p_i represents a distinct professional perspective. For a given query q , we compute its similarity to each perspective vector p_i , yielding a weight w_i that quantifies the alignment between q and p_i . To exclude weak alignments, weights are filtered by a threshold θ , which is determined during the offline stage by analyzing the statistical distribution of weights to ensure a consistent balance between relevance and coverage.

The weight w_i is computed as follows:

$$w_i = \begin{cases} \text{Similarity}(q, p_i) & \text{if } \text{Similarity}(q, p_i) > \theta \\ 0 & \text{otherwise,} \end{cases}$$

where $\text{Similarity}(q, p_i)$ represents the cosine similarity between the query q and perspective p_i . The threshold θ is chosen to balance the inclusion of relevant perspectives with the exclusion of noisy matches, typically optimized using a development set.

The resulting weights form a *Perspective Vector* V_q , represented as:

$$V_q = [w_1 \quad w_2 \quad \dots \quad w_n],$$

where the i -th element w_i reflects the relevance of perspective p_i to the query q .

Query Rewriting: Using the Perspective Vector V_q , the system rewrites the query q for each perspective p_i with a non-zero weight w_i . A large-scale language model, referred to as the **rewriter**,

generates rewritten content C_i for each such perspective:

$$C_i = \text{Rewriter}(q, p_i) \quad \forall p_i \text{ with } w_i \neq 0, \quad (5)$$

where q is the original query, p_i is the specific perspective, and w_i represents its weight. The rewriting process leverages the contextual understanding of the rewriter model to produce content tailored to the unique nuances of each perspective.

Contextual Document Retrieval: For each rewritten query C_i , the system retrieves a set of documents from the corpus using similarity-based search. The retrieved set R_i for perspective p_i is defined as:

$$R_i = \left\{ d_{ij} \left| \begin{array}{l} d_{ij} \text{ ranks among the top } k \text{ by} \\ \text{Similarity}(d_{ij}, C_i), \forall j \in \{1, \dots, k\} \end{array} \right. \right\}, \quad (6)$$

where d_{ij} represents the j -th document in R_i , and $\text{Similarity}(d_{ij}, C_i)$ is the cosine similarity between document d_{ij} and rewritten query C_i . k is a hyperparameter that defines the retrieval depth, dynamically set based on the user’s requirements for the desired scope and granularity of the retrieved results.

The retrieval process produces a collection of document sets $\{R_1, R_2, \dots, R_n\}$, each tailored to a specific professional perspective. This ensures the query is contextually expanded and aligned with the multifaceted aspects of the domain, significantly enhancing the relevance and specificity of the retrieved information.

3.4 Retrieval Augmentation and Final Inference

After assembling the retrieved document sets $\{R_1, R_2, \dots, R_n\}$, the system refines these results to enhance relevance and utility. This process is guided by the original *Perspective Vector* V_q , ensuring that the final output reflects the multi-perspective nature of the query.

Re-ranking Documents: The system recalculates the relevance of each document d_{ij} in the retrieved sets based on its alignment with the rewritten query C_i and the corresponding perspective weight w_i . The relevance score $\text{Rel}(d_{ij})$ is defined as:

$$\text{Rel}(d_{ij}) = \frac{\text{Similarity}(d_{ij}, C_i)}{w_i}, \quad (7)$$

where $\text{Similarity}(d_{ij}, C_i)$ measures the cosine similarity between the document d_{ij} and the rewritten

query C_i , and w_i represents the weight of perspective p_i in V_q . This formula ensures that documents strongly aligned with both the rewritten query and the perspective are prioritized.

Structured Prompt Generation: The re-ranked documents are integrated into a structured prompt $\mathcal{P}$, combining the original query q with the restructured document content:

$$\mathcal{P} = q \bowtie \bigoplus_{i=1}^n \bigoplus_{j=1}^{k_i} d'_{ij}, \quad (8)$$

where d'_{ij} are the re-ranked documents, $\bowtie$ represents concatenation, and $\bigoplus$ denotes the combination of results across all perspectives. The prompt ensures that all relevant perspectives are represented, with emphasis on their respective importance as determined by V_q .

Final Inference: The structured prompt $\mathcal{P}$ is fed into a large-scale reader model $\mathcal{M}$, such as a pretrained transformer-based language model, to generate the final response. The reader model processes the combined information to produce a nuanced, multi-perspective answer that is contextually aligned with the original query’s complexities.

This retrieval augmentation framework provides a comprehensive analysis of the query, combining diverse perspectives while ensuring specificity and contextual richness. The result is a tailored response that effectively addresses the complexities of specialized domains.

4 Experiments

4.1 Retrieval Task

4.1.1 Dataset and Experiment Setup

For the retrieval task, the MVRAG framework aims to retrieve diverse and comprehensive documents to better support downstream generation tasks, differing from traditional methods that focus on precision. Despite this distinction, experimental results show that MVRAG performs on par with or even surpasses traditional retrieval methods.

The evaluation was conducted using datasets from four domains: **e-commerce**, **medical**, **entertainment video**, and **legal**. The first three domains were sourced from Alibaba’s *Multi-CPR* (Long et al., 2022), containing approximately one million records and 1000 queries per domain. For the legal domain, the *LeCaRDv2* dataset (Li et al., 2023), a benchmark dataset comprising 800 queries and

Method	E-commerce			Medical			Entertainment			Legal		
	R@5	R@10	MRR@10	R@5	R@10	MRR@10	R@5	R@10	MRR@10	R@5	R@10	MRR@10
Traditional Models												
BM25	20.50%	26.00%	0.1575	31.60%	35.20%	0.2695	31.50%	40.10%	0.2198	15.12%	24.81%	0.1473
QLD	21.70%	24.70%	0.1629	36.10%	42.20%	0.2385	25.20%	26.10%	0.2032	14.46%	25.38%	0.1389
Dense Retrieval Models												
Bert	2.90%	4.30%	0.0257	1.60%	2.30%	0.0167	4.70%	6.00%	0.0403	4.37%	6.32%	0.0592
BGE-Large	45.60%	54.40%	0.3518	55.10%	59.60%	0.4885	43.10%	53.80%	0.3106	11.10%	18.30%	0.0913
BGE-M3	40.30%	48.30%	0.2977	53.60%	57.50%	0.4694	31.80%	42.10%	0.2366	12.09%	20.63%	0.1153
Dense Retrieval Models with MVRAG												
Bert	2.90%	5.80%	0.0307	9.00%	11.70%	0.0352	4.70%	7.20%	0.0390	2.23%	8.34%	0.0351
BGE-Large	48.80%	66.80%	0.3314	51.60%	59.40%	0.4550	45.90%	48.40%	0.2183	13.30%	24.95%	0.1087
BGE-M3	40.10%	56.10%	0.2910	60.40%	69.40%	0.5123	37.40%	42.40%	0.2413	16.09%	29.18%	0.1525

Table 1: Performance of retrieval methods across datasets. Metrics include recall (R) at 5 and 10 documents and mean reciprocal rank (MRR) at 10.

55,192 criminal case document candidates, was used.

To ensure a fair comparison, the baselines included traditional retrieval models BM25 and QLD, implemented with *Pyserini*⁹ default settings, and dense retrieval models such as **bert-base-chinese**¹⁰, **bge-large-zh-v1.5**¹¹, and **bge-m3**¹², configured according to standard *Dense* framework practices. In MVRAG, the same dense retrieval models and *Dense*¹³ framework were employed alongside **DeepSeek**¹⁴ as the query rewriter, with the retrieval depth k for each perspective set to 10.

4.1.2 Results

The experimental results demonstrate that the MVRAG framework, designed to retrieve diverse and comprehensive documents to support downstream tasks, achieves competitive performance across various datasets. While minor declines are observed in certain metrics, the majority remain comparable to or surpass baseline methods, particularly in knowledge-intensive domains.

In the **medical domain**, Recall@10 improves significantly from 57.50% to 69.40%, with MRR@10 increasing from 0.4694 to 0.5123. Similarly, in the **legal domain**, Recall@10 increases from 20.63% to 29.18%, and MRR@10 improves from 0.1153 to 0.1525. In less complex domains such as **e-commerce** and **entertainment video**, MVRAG consistently maintains or slightly improves performance, achieving a Recall@10 of

66.80% in e-commerce.

These findings underscore MVRAG’s robust retrieval capabilities across a wide range of datasets, even though its primary objective is to enhance the diversity and comprehensiveness of retrieved documents in knowledge-intensive scenarios.

4.2 Generation Task

4.2.1 Dataset and Experiment Setup

For the generation tasks, we demonstrated the improvements from the MVRAG framework using domain-specific question-answering tasks. In the legal domain, we used the *JEC-QA* dataset (Zhong et al., 2020), which contains 26,365 multiple-choice questions from the Chinese judicial examination, a challenging test for legal practitioners. The dataset was split into single-choice and multiple-choice questions for testing with LLMs, using the provided legal texts as the retrieval database. In the medical domain, we selected the *MED-QA* dataset¹⁵, consisting of 14,123 single-choice biomedical questions from the National Medical Licensing Examination.

We employed GLM-4-9B¹⁶ and Qwen-7B¹⁷ as rewriter and generator models for their strong performance in Chinese and efficient operation. Evaluation was based on accuracy for single-choice questions, and exact match accuracy and correct answer percentage for multiple-choice questions. For subjective questions, we used official standard answers and an automatic scoring system based on keyword matching. The retrieval depth for other models was set to 8, while for MVRAG, the retrieval depth k for each perspective was set to 2. This adjust-

⁹<https://github.com/castorini/pyserini>

¹⁰<https://huggingface.co/google-bert/bert-base-chinese>

¹¹<http://huggingface.co/BAAI/bge-large-zh-v1.5>

¹²<https://huggingface.co/BAAI/bge-m3>

¹³<https://github.com/luyug/Dense>

¹⁴<https://www.deepseek.com/>

¹⁵https://huggingface.co/datasets/bigbio/med_qa

¹⁶<https://github.com/THUDM/GLM-4>.

¹⁷<https://modelscope.cn/models/qwen/Qwen-7B>.

Model	Method	MS	LS	LM (Exact)	LM (Percentage)	LSub
GLM-4-9B	Baseline	68.72%	66.42%	5.49%	14.57%	18.5%
	RAG	74.39%	69.41%	7.03%	18.43%	19.35%
	Query2doc	78.31%	73.10%	13.93%	31.33%	22.94%
	MultiQuery	75.02%	70.43%	7.00%	23.58%	21.38%
	MVRAG	80.20%	72.50%	22.91%	48.42%	29.64%
Qwen-7B	Baseline	66.50%	72.13%	5.00%	14.57%	14.91%
	RAG	70.21%	75.76%	5.56%	18.89%	16.12%
	Query2doc	79.01%	76.94%	12.30%	35.12%	19.35%
	MultiQuery	77.96%	77.02%	13.14%	35.19%	20.72%
	MVRAG	78.83%	77.89%	26.26%	50.59%	30.69%

Table 2: Performance on legal and medical generation tasks. The abbreviations are as follows: MS (Medicine Single), LS (Legal Single), LM (Exact) (Exact match accuracy for Legal Multi-choice), LM (Percentage) (Percentage of correct answers for Legal Multi-choice), and LSub (Legal Subjective QA).

ment ensures fairness, as MVRAG extracts four perspectives during the offline perspective extraction phase. Our experiments were conducted in an environment equipped with a single A100 GPU.

4.2.2 Results

The results, shown in Table 2, demonstrate that MVRAG excels across tasks, particularly in complex scenarios. While single-choice questions showed relatively smaller gains due to well-defined answers, MVRAG still improved accuracy to 72.50% on GLM-4-9B and 77.89% on Qwen-7B. However, the most notable improvements were seen in multiple-choice and subjective tasks, where MVRAG’s multi-view query rewriting and feature extraction significantly enhanced model performance. For legal multiple-choice tasks, MVRAG boosted accuracy to 22.91% on GLM-4-9B and 26.26% on Qwen-7B, nearly tripling the baseline performance.

MVRAG’s strength lies in its ability to integrate domain-specific perspectives, enabling deeper analysis of complex legal and medical scenarios. By capturing multidimensional aspects like dispute focus in legal cases or medical history in diagnoses, MVRAG delivers more precise and contextually rich results. This multi-view approach explains why MVRAG consistently outperforms other methods in tasks requiring complex reasoning.

In summary, MVRAG demonstrates value in simpler tasks but shows transformative improvements in complex, knowledge-intensive queries, enhancing retrieval precision and inference quality in domains like law and medicine

4.3 Comparative Experiments on Professional Perspective Extraction

We compared traditional machine learning and LLM-based methods for **professional perspective**

extraction on a curated legal dataset. As detailed in Appendix A, the LLM-based method achieves better performance but at significantly higher computational costs. The machine learning approach, with its simplicity, efficiency, and lower resource demands, is better suited for practical deployment within the MVRAG framework.

4.4 Ablation Experiments

We performed ablation experiments on the legal and medical datasets to evaluate the effect of perspective selection on retrieval. In both legal and medical tasks, we compared the full model with versions omitting one perspective. The degradation in retrieval performance after removing specific perspectives demonstrates the effectiveness of our framework. See Appendix B for details.

5 Conclusion

In this paper, we propose a multi-perspective approach to Retrieval-Augmented Generation tailored for knowledge-dense domains, aiming to incorporate domain-specific insights missing from existing methods and enhance the reliability and interpretability of retrieval outcomes. By employing professional perspectives extraction, intent recognition, query rewriting, and document re-ranking, We have significantly improved generation performance while maintaining retrieval accuracy in fields such as law and medicine. Through experiments, we demonstrate the impact of integrating multi-perspective information, laying the groundwork for future incorporation of machine learning techniques into RAG systems. Our approach introduces advanced feature extraction techniques and unleashes the potential of multi-view information, accelerating the application of LLMs in knowledge-dense fields.

Limitations

While the MVRAG framework achieves notable improvements, it is not without limitations. Firstly, the use of domain-specific corpora and learned perspectives risks embedding and perpetuating biases present in the data, such as those related to race, gender, or other social constructs. These biases could influence retrieval outcomes, reinforcing existing prejudices and undermining fairness in sensitive applications like law or medicine. Secondly, the lack of tailored evaluation metrics for multi-view systems limits the ability to rigorously measure their effectiveness, particularly in assessing the diversity and equity of retrieved perspectives. Future research should address these issues by incorporating bias mitigation strategies in data pre-processing and perspective modeling, alongside developing comprehensive metrics to evaluate both fairness and retrieval quality.

References

- Jiaxi Cui, Zongjian Li, Yang Yan, Bohua Chen, and Li Yuan. 2023. Chatlaw: Open-source legal large language model with integrated external knowledge bases. *arXiv preprint arXiv:2306.16092*.
- Shahul Es, Jithin James, Luis Espinosa-Anke, and Steven Schockaert. 2023. Ragas: Automated evaluation of retrieval augmented generation. *arXiv preprint arXiv:2309.15217*.
- Luyu Gao, Xueguang Ma, Jimmy Lin, and Jamie Callan. 2022. Precise zero-shot dense retrieval without relevance labels, 2022. URL <https://arxiv.org/abs/2212.10496>.
- Yunfan Gao, Yun Xiong, Xinyu Gao, Kangxiang Jia, Jinliu Pan, Yuxi Bi, Yi Dai, Jiawei Sun, and Haofen Wang. 2023. Retrieval-augmented generation for large language models: A survey. *arXiv preprint arXiv:2312.10997*.
- Kelvin Guu, Kenton Lee, Zora Tung, Panupong Pasupat, and Mingwei Chang. 2020. Retrieval augmented language model pre-training. In *International conference on machine learning*, pages 3929–3938. PMLR.
- Siqing Huo, Negar Arabzadeh, and Charles LA Clarke. 2023. Retrieving supporting evidence for llms generated answers. *arXiv preprint arXiv:2306.13781*.
- Gautier Izacard, Patrick Lewis, Maria Lomeli, Lucas Hosseini, Fabio Petroni, Timo Schick, Jane Dwivedi-Yu, Armand Joulin, Sebastian Riedel, and Edouard Grave. 2022. Atlas: Few-shot learning with retrieval augmented language models. *arXiv preprint arXiv:2208.03299*.
- Omar Khattab, Christopher Potts, and Matei Zaharia. 2021. Relevance-guided supervision for openqa with colbert. *Transactions of the association for computational linguistics*, 9:929–944.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, et al. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in Neural Information Processing Systems*, 33:9459–9474.
- Haitao Li, Yunqiu Shao, Yueyue Wu, Qingyao Ai, Yixiao Ma, and Yiqun Liu. 2023. Lecardv2: A large-scale chinese legal case retrieval dataset. *arXiv preprint arXiv:2310.17609*.
- Dingkun Long, Qiong Gao, Kuan Zou, Guangwei Xu, Pengjun Xie, Ruijie Guo, Jian Xu, Guanjin Jiang, Luxi Xing, and Ping Yang. 2022. Multi-cpr: A multi domain chinese dataset for passage retrieval. In *SI-GIR*, pages 3046–3056. ACM.
- Xinbei Ma, Yeyun Gong, Pengcheng He, Hai Zhao, and Nan Duan. 2023. Query rewriting for retrieval-augmented large language models. *arXiv preprint arXiv:2305.14283*.
- Fabian Pedregosa, Gaël Varoquaux, Alexandre Gramfort, Vincent Michel, Bertrand Thirion, Olivier Grisel, Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vincent Dubourg, Jake Vanderplas, Alexandre Passos, David Cournapeau, Matthieu Brucher, Matthieu Perrot, and Edouard Duchesnay. 2011. Scikit-learn: Machine learning in python. *Journal of Machine Learning Research*, 12:2825–2830.
- Ori Ram, Yoav Levine, Itay Dalmedigos, Dor Muhlgay, Amnon Shashua, Kevin Leyton-Brown, and Yoav Shoham. 2023. In-context retrieval-augmented language models. *Transactions of the Association for Computational Linguistics*, 11:1316–1331.
- Kurt Shuster, Spencer Poff, Moya Chen, Douwe Kiela, and Jason Weston. 2021. Retrieval augmentation reduces hallucination in conversation. *arXiv preprint arXiv:2104.07567*.
- Liang Wang, Nan Yang, and Furu Wei. 2023. Query2doc: Query expansion with large language models. *arXiv preprint arXiv:2303.07678*.
- Yunjia Xi, Weiwen Liu, Jianghao Lin, Jieming Zhu, Bo Chen, Ruiming Tang, Weinan Zhang, Rui Zhang, and Yong Yu. 2023. Towards open-world recommendation with knowledge augmentation from large language models. *arXiv preprint arXiv:2306.10933*.
- Jing Yao, Wei Xu, Jianxun Lian, Xiting Wang, Xiaoyuan Yi, and Xing Xie. 2023. Knowledge plugins: Enhancing large language models for domain-specific recommendations. *arXiv preprint arXiv:2311.10779*.
- Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, et al. 2023. A

677 survey of large language models. *arXiv preprint*
678 *arXiv:2303.18223*.

679 Haoxi Zhong, Chaojun Xiao, Cunchao Tu, Tianyang
680 Zhang, Zhiyuan Liu, and Maosong Sun. 2020. Jec-
681 qa: a legal-domain question answering dataset. In
682 *Proceedings of the AAAI conference on artificial in-*
683 *telligence*, volume 34, pages 9701–9708.

684 Yutao Zhu, Huaying Yuan, Shuting Wang, Jiongnan
685 Liu, Wenhan Liu, Chenlong Deng, Zhicheng Dou,
686 and Ji-Rong Wen. 2023. Large language models
687 for information retrieval: A survey. *arXiv preprint*
688 *arXiv:2308.07107*.

Appendix A Comparative Experiments on Professional Perspective Extraction

To compare traditional machine learning and LLM-based methods for professional perspective extraction, we conducted experiments on a curated legal dataset using GLM4-9B as the baseline model. The results, summarized in Table 3, indicate that while the LLM-based method often generates more nuanced perspectives, it requires significantly higher computational time and resources.

For instance, the machine learning approach completed extraction within minutes, whereas the LLM-based method required hours, highlighting its resource-intensive nature. Despite this, the performance of the machine learning approach was competitive, offering simplicity, efficiency, and ease of integration. These advantages make it particularly suitable for practical deployment in the MVRAG framework, especially in resource-constrained scenarios.

Overall, the findings underscore the trade-offs between the two methods: the LLM-based approach excels in generating richer perspectives but is less practical for large-scale or time-sensitive applications. In contrast, the machine learning method balances performance and efficiency, aligning well with the objectives of the MVRAG framework.

Method	Time Taken	Extracted Perspectives	Legal Multiple	Legal Subjective
ML	1.5 minutes	<i>Basic Fact, Focus of Dispute</i> <i>Application of Law, Penalty, Criminal History</i>	48.42%	29.64%
LLM	2 hours	<i>Essential Case Facts, Primary Legal Issue</i> <i>Legal Precedent Interpretation, Criminal History, Sanction</i>	39.98%	26.45%

Table 3: Comparison of Machine Learning (ML) and LLM-based methods (LLM) in professional perspective extraction across various metrics and legal perspectives.

Appendix B Ablation Study

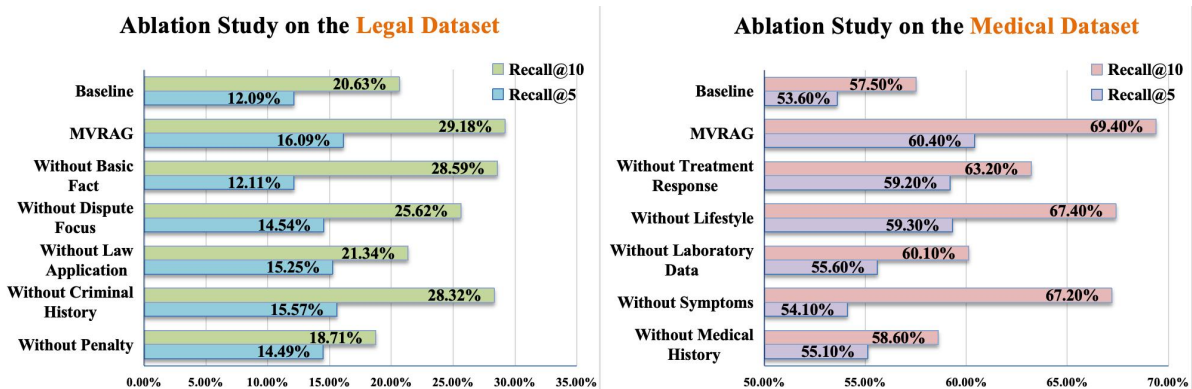


Figure 5: Ablation Study on the impact of perspective selection strategies in our framework on Medical and Legal datasets. In the legal domain, the chart shows Recall@5 and Recall@10 after excluding each perspective: *Basic Fact*, *Focus of Dispute*, *Application of Law*, *Penalty*, *Criminal History*. For the medical domain, it displays the effects of removing *Medical History*, *Symptoms*, *Laboratory Data*, *Treatment Response*, *Lifestyle*. Each bar indicates the performance impact versus the full baseline and direct retrieval.

To evaluate the impact of individual perspectives on retrieval performance, we conducted ablation experiments on the medical and legal datasets. Each dataset employs a tailored multi-perspective retrieval framework, comprising five key perspectives. In the medical domain, the perspectives are *Medical History*, *Symptoms*, *Laboratory Data*, *Lifestyle*, and *Treatment Response*. For the legal domain, the perspectives include *Fundamental Facts*, *Focus of Dispute*, *Application of Law*, *Criminal History*, and *Penalty*. We compared the full model using all perspectives against configurations where one perspective was omitted at a time.

B.1 Medical Domain Results

In the medical dataset, the full model achieved the highest retrieval performance with a *Recall@5* of 60.4% and *Recall@10* of 69.4%. Removing individual perspectives revealed their relative contributions to retrieval accuracy:

- **Medical History:** Excluding this perspective caused the most significant performance drop, with *Recall@5* declining to 55.1% and *Recall@10* to 58.6%, indicating its critical role in capturing relevant cases.
- **Symptoms:** The absence of this perspective reduced *Recall@5* to 54.1%, while *Recall@10* remained relatively high at 67.2%.
- **Laboratory Data:** Removing this perspective resulted in *Recall@5* of 55.6% and *Recall@10* of 60.1%.
- **Lifestyle:** Excluding this perspective had the least impact, maintaining *Recall@5* at 59.3% and *Recall@10* at 67.4%.
- **Treatment Response:** Omission of this perspective led to a moderate reduction in performance, with *Recall@5* dropping to 59.2% and *Recall@10* to 63.2%.

B.2 Legal Domain Results

For the legal dataset, the full multi-perspective framework achieved a *Recall@5* of 15.57% and a *Recall@10* of 28.59%. The ablation study revealed the following:

- **Fundamental Facts:** Excluding this perspective caused the most substantial performance drop, with *Recall@5* declining to 12.11% and *Recall@10* to 18.71%, underscoring its foundational role in the retrieval process.
- **Focus of Dispute:** The omission of this perspective led to *Recall@5* of 14.54% and *Recall@10* of 25.62%, showing its significant but lesser contribution compared to *Fundamental Facts*.
- **Application of Law:** Removing this perspective resulted in *Recall@5* of 15.25% and *Recall@10* of 21.34%.
- **Criminal History:** Excluding this perspective caused a moderate performance reduction, with *Recall@5* decreasing to 15.57% and *Recall@10* to 28.32%.
- **Penalty:** The absence of this perspective had the least impact, reducing *Recall@5* slightly to 14.49% and *Recall@10* to 18.71%.

B.3 Discussion

The results demonstrate the effectiveness of the multi-perspective retrieval framework in both domains. In the medical dataset, perspectives such as *Medical History* play a pivotal role, while perspectives like *Lifestyle* show partial redundancy. Similarly, in the legal dataset, *Fundamental Facts* emerged as the most critical perspective, while *Penalty* contributed the least to overall performance. These findings highlight the nuanced contributions of each perspective, providing insights for optimizing retrieval strategies in domain-specific contexts.

Appendix C Methodological Details

In this appendix, we provide detailed implementation specifics of several key methodological steps in our main framework.

C.1 Multi-View Query Rewriting Using Zero-Shot In-Context Learning

To enable scalability and adaptability across diverse domains, we employ a zero-shot in-context learning approach for generating rewritten queries from multiple perspectives. This method empowers the language model to generalize its understanding of the task without relying on explicit examples, making it suitable for a wide range of applications such as medical, legal, and scientific domains. Below, we provide an example of the zero-shot prompt applied in the legal domain:

Prompt Example

You are a domain expert specializing in analyzing and rewriting queries to provide comprehensive information from different perspectives. Given the following input, please generate the corresponding pseudo-answer document.

Input Query:

In a theft case, the suspect was caught at the crime scene but claims he did not participate in the theft. How should this situation be handled?

List of Perspectives:

1. Basic Facts Perspective
2. Focus of Dispute Perspective
3. Legal Application Perspective
4. Punishment Perspective

Pseudo-Answer Document:

1. Basic Facts Perspective:
 - Pseudo-Answer: Describe the objective facts related to the case, such as the suspect being caught at the crime scene, the evidence collected, and any statements from eyewitnesses.
2. Focus of Dispute Perspective:
 - Pseudo-Answer: Explain the key points of contention, such as the suspect's claim of innocence versus the physical evidence and eyewitness accounts.
3. Legal Application Perspective:
 - Pseudo-Answer: Provide an analysis of relevant legal principles or statutes, such as the definition of theft under Article 264 of the Criminal Law and the requirement for conclusive evidence.
4. Punishment Perspective:
 - Pseudo-Answer: Suggest an appropriate penalty if guilt is established, referencing sentencing guidelines under Article 264 of the Criminal Law.

This prompt demonstrates the flexibility of zero-shot in-context learning, where the language model interprets the task requirements without explicit examples. While the above case pertains to the legal domain, the same framework can seamlessly adapt to other domains, such as generating multi-perspective queries for medical or scientific tasks. This adaptability ensures the broad applicability of the multi-view query rewriting approach across various knowledge-intensive fields.

C.2 Multi-View Query Answering Using Zero-Shot In-Context Learning

To answer legal multiple-choice questions using zero-shot in-context learning, the following prompt is used to guide an advanced language model in selecting the appropriate answers based solely on a given question and its contextual reference, without relying on an example.

Prompt Example

You are a legal professional tasked with answering the following legal exam multiple-choice question. Based on the provided question and the given reference context, select one or more correct answers without providing any analysis or reasoning.

Your Task

- Instruction: {instruction}
- Question and Options: {question}
- Search Context: {search_context}

Please note: the reference information might contain inaccuracies, so proceed with caution. Only select the most appropriate answer(s) based on the context provided.

This zero-shot prompt directs the model to generate an answer without needing a prior example, making it efficient for generating answers to new queries quickly, based on minimal input and context.

Appendix D A Typical Case Study in the Field of Medicine

To better showcase the effectiveness of our model, we conduct some case studies in several knowledge-dense domains, including *Biology*, *Geography*, *Literature*, *Political Science*, and *Physics*. An extensive collection of detailed case studies across a wide range of disciplines are provided in the Appendix E for further exploration.

Within this section, we delve into the medical domain, employing the *PMC-Patients* dataset as the retrieval database. As illustrated in Figure 2, we selected a case of **Huntington’s disease** from *PubMed* and transformed it into a query voiced by the patient, serving as our original query. In scenarios absent of our multi-view retrieval model, the retrieval system primarily fetched articles related to superficial symptoms such as emotional dysregulation and movement disorders. This surface-level matching, due to similar characterizations, erroneously led to articles on **Vitamin B12 deficiency**, consequently misdiagnosing the condition as a **Vitamin B12 deficiency**.

Conversely, our multi-view retrieval system initiates with intent recognition, prioritizing symptoms and medical history for query reformulation. The perspectives utilized by the framework are determined during an offline stage, where machine learning techniques such as Principal Component Analysis (PCA) and Non-negative Matrix Factorization (NMF) analyze the structural and semantic patterns of the domain-specific corpus. This automated process extracts professional perspectives, ensuring they are representative of the corpus’s inherent knowledge while being adaptable to various domains.

Through individual matching and re-ranking, the system yielded comparatively relevant search outcomes. Specifically, symptom keywords, refined with greater precision, directed the search towards articles on neurodegenerative diseases. Simultaneously, the emphasis on similar family medical history surfaced articles on genetic testing. With the aid of these references, the model adeptly identified the potential for **Huntington’s disease**.

This case study emphatically demonstrates the superiority of our multi-view retrieval model. By leveraging perspectives refined through offline corpus analysis and multi-view reformulation, the system adeptly navigated towards more pertinent and informative references, thereby significantly enhancing the accuracy of the final diagnostic inference.

Appendix E Case Study Across Different Fields

E.1 Biology

Domain	Domain Perspectives
Biology	Gene Expression Protein Function Species Evolution
Original Query	Rewritten Queries
Why do some species' pigment deposits change with environmental variations?	<i>Gene Expression Perspective:</i> "How do environmental factors affect gene expression in species leading to changes in pigment deposition?" <i>Protein Function Perspective:</i> "What impact does environmental change have on the function and activity of pigment proteins in species?" <i>Species Evolution Perspective:</i> "How do changes in pigment deposition reflect the adaptation and evolution of species to environmental changes?"
Direct Retrieval	Multi-View Retrieval
"Environmental Pollution and Animal Pigment Changes" "The Impact of Global Warming on Plant Colors" "The Effect of Light on Butterfly Wing Coloration"	<i>Gene Expression Perspective:</i> "Gene Expression Regulation and Pigment Accumulation under Environmental Stress" "Adaptive Gene Expression Changes in Response to Environmental Light Variations" <i>Protein Function Perspective:</i> "The Role of Pigment Proteins in Environmental Adaptation" "Functional Diversity of Pigment Proteins under Environmental Pressure" <i>Species Evolution Perspective:</i> "Molecular Mechanisms of Species Color Evolution Prompted by Environmental Changes" "Evolutionary Adaptations in Pigmentation: Evidence from Environmental Interactions"
Inference	Multi-View Inference

Based on our findings, it appears that changes in species' pigment deposits are primarily due to direct environmental impacts, such as pollution, global warming, and variations in light exposure. These external factors seem to straightforwardly affect the coloration without involving complex biological processes.	Upon reviewing comprehensive articles across multiple perspectives, we can conclude that the variations in pigment deposits among species in response to environmental changes are not merely direct reactions. Instead, these changes are the result of intricate biological mechanisms. Environmental factors influence gene expression and the functionality of pigment proteins, which in turn contribute to the species' evolutionary adaptation to their habitats. Specifically, temperature and light variations play crucial roles in modulating gene expression related to pigment deposition, highlighting the adaptive and evolutionary significance of these pigment changes.
--	---

798

E.2 Geography

Domain	Domain Perspectives
Geography	Climate Change Urban Development Biodiversity
Original Query	Rewritten Queries
How does the expansion of urban areas affect local climates and ecosystems?	<i>Climate Change Perspective:</i> "In what ways does urban expansion contribute to changes in local climate patterns?" <i>Urban Development Perspective:</i> "How does the planning and growth of urban areas influence their surrounding ecosystems and climate?" <i>Biodiversity Perspective:</i> "What impact does urban sprawl have on the biodiversity of adjacent natural habitats?"
Direct Retrieval	Multi-View Retrieval

"The Heat Island Effect in Mega Cities" "Urbanization and Its Impact on Rainfall Variability" "Effects of Concrete Surfaces on Urban Temperatures"	<i>Climate Change Perspective:</i> "Urban Expansion and Its Role in Altering Regional Climate Systems" "Greenhouse Gas Emissions from Urban Centers: A Climate Change Perspective" <i>Urban Development Perspective:</i> "Sustainable Urban Planning: Balancing Growth and Environmental Preservation" "The Influence of Urban Landscape Design on Local Ecosystems" <i>Biodiversity Perspective:</i> "Urbanization's Toll on Local Wildlife: Case Studies of Habitat Fragmentation" "Conservation Strategies for Mitigating Urbanization Effects on Biodiversity"
Inference The research primarily points to urban areas contributing to higher temperatures through the heat island effect and altering rainfall patterns due to extensive concrete surfaces. It suggests that the primary impact of urban expansion is a direct increase in local temperatures and changes in precipitation.	Multi-View Inference After analyzing data from various perspectives, it is evident that urban expansion affects local climates and ecosystems in multiple interconnected ways. Urban growth leads to increased greenhouse gas emissions, significantly influencing local climate patterns beyond just temperature increases, such as altering precipitation and wind patterns. From an urban development viewpoint, the layout and planning of cities can either harm or help preserve local ecosystems, highlighting the importance of sustainable development to mitigate negative impacts. Furthermore, urban sprawl severely affects biodiversity, leading to habitat fragmentation and loss. Therefore, the expansion of urban areas requires careful consideration of its broader environmental impacts, necessitating strategies that prioritize sustainability and biodiversity conservation.

E.3 Literature

799

Domain	Domain Perspectives
Literature	Narrative Techniques Historical Context Cultural Impact
Original Query	Rewritten Queries

How do modern novels reflect contemporary societal issues?	<i>Narrative Techniques Perspective:</i> "What narrative techniques are modern novelists using to explore and reflect contemporary societal issues?" <i>Historical Context Perspective:</i> "How does the historical context of the early 21st century influence the themes and subjects of modern novels?" <i>Cultural Impact Perspective:</i> "In what ways do modern novels influence and reflect the cultural attitudes and social issues of our time?"
Direct Retrieval	Multi-View Retrieval
"The Rise of the Digital Novel: Technology in Modern Literature" "Post-Modernism and Its Influence on 21st Century Literature" "The Evolution of Character Archetypes in Contemporary Fiction"	<i>Narrative Techniques Perspective:</i> "Stream of Consciousness and Nonlinear Narratives in Depicting Modern Complexities" "The Role of Metafiction in Critiquing Contemporary Societal Norms" <i>Historical Context Perspective:</i> "Post-9/11 Literature: Terrorism and Its Aftermath in Modern Novels" "Economic Crises and Their Reflections in 21st Century Fiction" <i>Cultural Impact Perspective:</i> "Literature as a Mirror to the #MeToo Movement: Narratives of Empowerment and Justice" "Climate Fiction: How Modern Novels Address Environmental Concerns"
Inference	Multi-View Inference

Based on the articles retrieved, it appears that modern novels primarily explore technological advancements and their implications, with a focus on the stylistic elements of post-modernism and the evolution of character archetypes. This suggests a significant emphasis on form and structure over content when reflecting contemporary societal issues.	Upon examining articles from diverse perspectives, it becomes clear that modern novels deeply engage with contemporary societal issues through various means. Through innovative narrative techniques, they offer nuanced explorations of complex issues like global terrorism, environmental crises, and social justice movements, such as #MeToo. The historical context of the early 21st century, marked by significant global events and economic turmoil, profoundly influences the thematic preoccupations of contemporary literature. Moreover, modern novels not only reflect but also actively shape cultural attitudes towards pressing social issues, demonstrating literature's power to influence societal change and public discourse. Therefore, modern novels serve as a crucial lens through which the multifaceted concerns and dynamics of contemporary society are examined and understood.
---	--

E.4 Political Science

800

Domain	Domain Perspectives
Political Science	Information Dissemination Voter Behavior Regulatory Policies
Original Query	Rewritten Queries
How do social media platforms influence political discourse and public opinion?	Information Dissemination Perspective: "What role do social media platforms play in the spread of political information and the formation of public opinion?" Voter Behavior Perspective: "How are social media platforms affecting voter behavior and electoral outcomes in democratic societies?" Regulatory Policies Perspective: "What regulatory measures are being implemented to ensure the integrity of political discourse on social media platforms?"
Direct Retrieval	Multi-View Retrieval

"The Rise of Social Media in Political Campaigns" "Social Media: The New Public Square for Political Discussion" "Echo Chambers and Filter Bubbles: The Polarization of Political Discourse on Social Media"	Information Dissemination Perspective: "Algorithmic Bias and News Feed Algorithms: Shaping Political Information on Social Media" "The Role of Social Media in Civic Engagement and Political Mobilization" Voter Behavior Perspective: "Social Media's Impact on Voter Turnout and Political Participation: A Global Perspective" "The Influence of Online Political Advertising on Voter Preferences and Decisions" Regulatory Policies Perspective: "Regulating Political Advertising on Social Media: Challenges and Approaches" "Social Media Governance: Balancing Free Speech with Political Integrity"
Inference Based on the directly retrieved articles, it might be concluded that social media primarily serves as a new platform for political campaigns, facilitating public political discussions and contributing to political polarization through echo chambers and filter bubbles. This perspective suggests a neutral to negative impact, focusing on the divisive aspects of social media's role in politics.	Multi-View Inference Integrating insights from articles across multiple perspectives reveals a nuanced understanding of social media's role in political discourse and public opinion. Social media platforms are crucial for disseminating political information, engaging citizens in political processes, and mobilizing voter participation. However, challenges such as algorithmic biases can shape political information in ways that may not always be transparent or equitable, potentially influencing voter behavior and preferences. Furthermore, the advent of online political advertising highlights the need for regulatory policies to ensure the integrity of political discourse, balancing the protection of free speech with the prevention of misinformation and the maintenance of political integrity. To address these complexities effectively, a comprehensive approach involving all stakeholders—platforms, policymakers, and the public—is essential for harnessing social media's potential as a force for democratic engagement while mitigating its risks.

E.5 Physics

Domain	Domain Perspectives
Physics	Quantum Mechanics Cryptography Technological Advancement
Original Query	Rewritten Queries

What are the implications of quantum computing for traditional encryption methods?	<i>Quantum Mechanics Perspective:</i> "How do principles of quantum mechanics underpin the operation of quantum computers and their impact on encryption?" <i>Cryptography Perspective:</i> "What challenges does quantum computing pose to current cryptographic algorithms and traditional encryption methods?" <i>Technological Advancement Perspective:</i> "Considering the advancements in quantum computing, what are the future prospects for secure communication technologies?"
Direct Retrieval	Multi-View Retrieval
"Introduction to Quantum Computing" "The Basics of Cryptography" "Quantum Computers and the Future of Encryption"	<i>Quantum Mechanics Perspective:</i> "Entanglement and Superposition: Pillars of Quantum Computing" "Quantum Algorithms: Redefining Computational Capabilities and Encryption" <i>Cryptography Perspective:</i> "Quantum Computing vs. RSA Encryption: A Cryptographic Analysis" "Post-Quantum Cryptography: Preparing for the Quantum Challenge" <i>Technological Advancement Perspective:</i> "The Race for Quantum Supremacy: Implications for Global Communication Security" "Innovations in Quantum-Resistant Encryption Methods for Future Technologies"
Inference	Multi-View Inference

From the directly retrieved articles, one might infer that quantum computing represents a futuristic advancement with the potential to enhance computational speed and efficiency, impacting encryption in undefined ways. This perspective might suggest a gradual and manageable transition from traditional to quantum-resistant encryption methods, underestimating the immediacy and severity of the challenge posed by quantum computing to current encryption standards.	Upon analyzing articles from a multi-perspective approach, it becomes evident that quantum computing introduces profound implications for traditional encryption methods. Principles of quantum mechanics, such as entanglement and superposition, enable quantum computers to potentially break current cryptographic algorithms, including RSA, in a fraction of the time taken by classical computers. This reality necessitates urgent development in post-quantum cryptography to safeguard secure communication technologies against quantum attacks. The ongoing advancements in quantum computing not only challenge existing cryptographic practices but also push the boundary of technological innovation, driving the creation of quantum-resistant encryption methods. Thus, the transition to quantum-safe encryption is not merely an incremental upgrade but a critical evolution required to maintain the confidentiality and integrity of digital communications in the quantum era.
--	---