

Mycroft - Retrieval Augmented Generation for SDK Documentation

Diego Costa, Gabriel Matos, Gilson Russo, Leon Barroso, and Erick Bezerra

SIDIA
Manaus - AM, Brazil

Abstract. Information retrieval plays an important role in everyday tasks, especially when it comes to documentation. Retrieving information about private documentation used to build other software is very challenging due to its absence on the internet, meaning there is no information about it beyond its own documentation. Due to concerns about confidential data, using external proprietary systems is prohibited. Motivated by this, in this study, we present Mycroft, a retrieval system that leverages the Retrieval Augmented Generation technique to find a feasible approach that improves search and information retrieval requested by users about the documentation. To implement this system, a dataset of questions and answers about the documentation was generated for evaluation. The system was developed on-premise using open-source Large Language Models and evaluated using Natural Language Processing metrics and human evaluation to validate the generated answers. Following an extensive evaluation of the results, the proposed retrieval system demonstrated satisfactory performance in addressing user queries and achieved favorable outcomes in human evaluation, indicating its utility.

Keywords: Retrieval Augmented Generation, Large Language Models, Software Documentation.

1 Introduction

Information Retrieval has developed an important role in day-to-day tasks. It has become part of daily routines for most people due to the need for rapid information and decision-making [1, 2]. It is present in many forms, such as internet browsing, virtual assistants, and chatbots [1]. With the advance of deep learning algorithms, information retrieval systems have been significantly improved [3], especially with the rise of Large Language Models (LLMs), which have been revealed as suitable for question-answering applications [4].

With the capacity of LLMs in question answering, it is possible to reduce both time and effort in information retrieval [3]. However, using cloud computing tools and commercial LLM raises security concerns due to the possibility of data leaks. As an alternative, open-source LLMs can be deployed on-premise server [3].

Large Language Models require prior knowledge acquired from training about a subject to answer a query. When unaware of the subject, it hallucinates [5]. An approach for building LLM applications without training them is the Retrieval-Augmented Generation (RAG) [3], which addresses this limitation by using an external source of information as context for response generation [6, 7].

RAG systems are considered a good choice for handling private information, as they leverage the LLMs' Natural Language Processing (NLP) capabilities without requiring the computational power needed to train an LLM model [3]. RAG systems are being implemented in various fields such as health, law, software, biology, and finance [8–12].

These systems are built on proprietary organizational documents that were not part of the LLM training. These documents are extensive and vital for private and governmental organizations, sometimes making it a time-consuming task to search and consult this information in daily routines [3].

In this work, Mycroft (a RAG system) is proposed to answer user questions about a private Software Development Kit (SDK) used internally in the organization. For this

purpose, the documentation is used as a source of information by indexing it in a vector store database and retrieve it for the LLM context to answer the user’s query.

This paper is organized as follows: Section 2 presents the background and related works, listing similar contributions and results achieved. Section 3 outlines the methodology, describing the approaches used to achieve the results. Section 4 reports the results, Section 5 presents the lessons learned and threats to validity, while Section 6 concludes by discussing the experimental outcomes.

2 Background

LLMs are trained with massive databases from different areas of knowledge, enabling them to answer a wide range of questions across various subjects [13]. However, when it comes to private data, they have a high probability of hallucinating in their responses [5], as they have never encountered such data. One of the commonly used approaches to address this issue is the RAG [6].

The RAG approach consists of a pipeline that first breaks down documents into small pieces, then indexes these pieces in a vector database. This database is consulted when a query about related documentation is made. The information relevant to the query is retrieved and passed to the LLM as context to answer the user’s question [6]. To enhance the LLM response, various strategies are employed in the RAG pipeline. These include breaking down documents into different sizes, indexing information as metadata in the database, utilizing entity graphs, and leveraging other LLMs to retrieve more refined data for the LLM context. All these strategies are detailed in the Related Works subsection.

The RAG has been applied in various domain-specific areas and has demonstrated significant efficiency with adjustments in the pipeline for specific tasks in knowledge domains. Some authors modify the chunking strategy, the embedding model, the LLM’s prompt, and the LLM model, as seen in the papers studied in the survey [14]. One advantage of RAG is that it is easier to keep the LLM model response up-to-date, as it only requires updating the external source of information. Consequently, the RAG system will use this updated version to query and answer queries [15].

To evaluate the generated answers against ground truth, we employed the metrics BLEU [16] and ROUGE [17], which are widely adopted in related works cited in this paper. Human evaluation was used to assess the proposed RAG answers by verifying their generated responses based on the evaluation method outlined in paper [10].

Additionally, we evaluated semantic similarity using the RAGAS¹ framework and BERTScore [18]. These two metrics involve converting candidate and reference texts into numeric vector representations and measuring their semantic proximity using cosine similarity to calculate the angle between the two vectors [19].

2.1 Related Works

Recently, different types of RAG architectures have been applied in various domains, aiming to achieve the best balance between answer generation and computational resources. For this purpose, small open-source LLMs with 1, 3, 7, and 8 billion parameters have been widely used across a variety of research fields, as described in the works [20, 9, 10, 3, 12] where different small open-source LLMs are employed to generate answers. These models differ in the specific domain knowledge applied by RAG and the pipeline used for information retrieval and answer generation. To achieve better results, various adjustments are made.

¹ <https://docs.ragas.io/>

In the work [20] the authors used 12 different metrics to evaluate the LLM models' responses. Mistral:7B achieved the best score in 10 of them, although the authors used Mistral:7B to generate answers and questions for experiments, which may have biased the results.

The work [9] utilized Llama3.2:3B for chunk relevance checking and query refinement and employed Mixtral8x7B, Llama3.1:8B, and Gemma2:9B for answer generation. The authors compared the results using human evaluation and cosine similarity. Llama3.1:8B secured the best scores in both evaluation types.

The study [10] the LLM models Gemma2:2B, Llama3.2:1B and Phi3-mini:3.8B were utilized. The authors evaluated the context retrieval phase by calculating context recall using the RAGAS framework and an LLM as a judge, comparing retrieved chunks with the ground truth reference. The chosen LLM for this task was Llama3:8B, and to evaluate generated answers, they used Natural Language Processing metrics such as Bilingual Evaluation Understudy (BLEU), Recall-Oriented Understudy for Gisting Evaluation (ROUGE), and Metric for Evaluation of Translation with Explicit Ordering (METEOR). The best embedding model was stella.v5, and the best LLM for generating answers was Phi3-mini.

The work [3] proposed an entity tree, which is searched alongside relevant document retrieval from the database. If a user query mentions an entity, it is searched in the tree and added to the LLM context along with the retrieved documents. They experimented with the Llama2:7B base model and a fine-tuned version answered 21 questions correctly. However, the author also used Llama2:7B to generate questions and answers for validation.

The work [12] focused on how document chunks were split. They utilized a sentence transformer trained on over 256M questions and answers to split documents into sizes of 128, 256, and 512 tokens. Texts smaller than 2048 characters were concatenated until this limit was reached or a title or table was encountered in the document. They evaluated retrieval accuracy using ROUGE and BLEU metrics by comparing retrieved chunks with ground truth paragraphs. A chunk size of 512 yielded better results for answer generation, as it captures more context. However, in some cases, it adds irrelevant information to the context and misses important details.

3 Mycroft - Retrieval Augmented Generation for SDK Documentation

Coding scripts with an internally developed SDK tool is not always easy, even with its documentation, it can take some time to find answers to the questions that arise during the development phase.

In this regard, software engineers frequently consult technical documentation not only to create new scripts but also to maintain existing scripts with new features that are released and upgrade deprecations.

One concern is about confidential data [21], using external proprietary LLMs such as ChatGPT², Grok³ and Claude⁴ requires sending the data to external servers, which may facilitate data leaks [22].

To mitigate this issue, open source LLMs on-premise can be used [23]. Another concern is about computational capacity, as LLMs increase in size they also demand more computational resources. Therefore, small models are more suitable for our resources.

² <https://openai.com/>

³ <https://x.ai/grok>

⁴ <https://claude.ai/>

This section describes the methodology used to develop a RAG pipeline for retrieving information about a private Python SDK documentation used internally within an organization. It outlines the creation of a dataset for test experiments and the RAG approach, which includes documentation pre-processing, embedding generation, and evaluation metrics.

3.1 Dataset Creation

To validate RAG efficiency, a validation dataset is required. To create this dataset, the SDK documentation was divided into 2000-character chunks and passed to the LLM Qwen2.5:14B [24] with a prompt to generate relevant and concise questions with ground truth answers for each chunk. Care was taken to exclude the LLM used for generating the questions and answers dataset to avoid bias.

The questions and answers then underwent a human evaluation process to review their correctness and relevance. Those deemed irrelevant or incorrect were enhanced using additional relevant chunks to provide a more complete context for the LLM to regenerate the question and answer.

A second human evaluation followed, ensuring no errors remained and classifying the questions into three difficulty levels: easy, medium, and hard. Easy-level questions had straightforward answers in the documentation, medium-level questions required combining two different chunks for a complete answer, and hard-level questions needed more than two relevant chunks and involved combining retrieved information with user-provided details. By the end of this phase, 302 questions with verified ground truth answers were generated: 29 hard, 84 medium, and 189 easy.

3.2 RAG methodology

The objective of this study is to facilitate the process of retrieving information about SDK documentation to assist software engineers in their development and maintenance tasks. Therefore, it explored the combination of different techniques to extract relevant information and used 4 LLMs to generate answers for domain-specific questions. We defined 3 research questions:

- RQ1: Which combination of documentation chunking and embedding model contributed to the most accurate responses?
- RQ2: Which Large Language Model provided the best responses?
- RQ3: How good are the responses according to users' opinion?

Documentation chunking. To enhance context information for generative LLMs, two approaches for document chunking were tested. The first step involved dividing and identifying text and Python code within each document section, then separating them into text and code sections.

One chunk approach used a text recursive splitter to break down the section's text into chunks of 500 characters with 100-character overlaps. Then, we used a Python recursive splitter with the same number of characters as the text chunks to segment the Python code, as illustrated in Figure 1.

The other approach involves directly using the documentation sections that contain the class, its description, and code as a single chunk. This method keeps the entire section's information together within one chunk, while another chunk contains information about a different section. This approach is illustrated in Figure 1.

Segmenting the documentation into chunks improves contextual relevancy and prevents the LLM context window from overflowing with large documents to process [25]. All chunks are less than 1300 tokens, and all LLMs used in the experiments have context windows up to 8k tokens. The objective of these two chunk approaches is to determine which type (smaller or large) has a greater impact on the LLM answer.

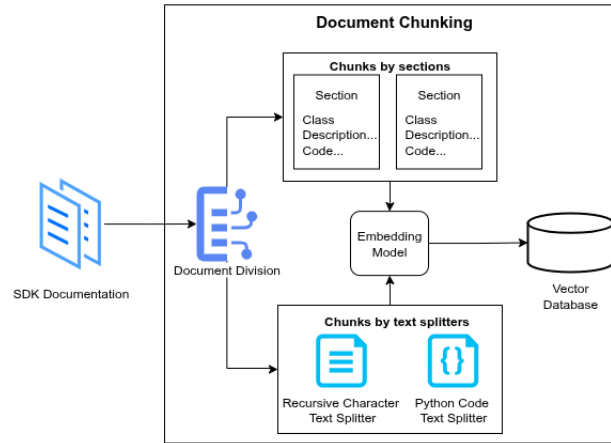


Fig. 1. Chunk strategies

Indexing. In this phase, the chunks extracted from the documentation undergo a process to transform them into numeric vector representations called vector embeddings, which are stored in a vector database [26]. Two embedding models, BGE-M3 [27] and nomic-embed-text-v1.5 [28], were selected to generate the vector embeddings for indexing in the Qdrant⁵ vector database. The documentation chunking and indexing phases are executed only once to store the documentation in the database. Consequently, this information is consulted during information retrieval.

Retrieval After the documentation is divided into chunks and indexed by Qdrant, the retrieval process is ready to start receiving queries about the SDK documentation. Qdrant will convert the question into a vector embedding chunk using the same embedding model, then it will search for similar vectors in the database using semantic similarity. By default, the search will retrieve 4 documents that will be passed to the generative LLM, which will use these documents as context to formulate an answer for the user query, as shown in Figure 2.

Generation During the generation phase, the LLM is prompted to use the retrieved documents from Qdrant and to answer the user's query based on this domain-specific retrieved context. The prompt guides the LLM to act as an assistant providing explanations on how to use a Python SDK according to the retrieved context.

If the retrieved context is insufficient to answer the query, the LLM is instructed to state that it lacks information, which prevents hallucinations. At the end of the prompt, a retrieved context field is included, populated with documents from the Qdrant database, and a user query field. When the LLM is called, these fields are filled with actual data.

⁵ <https://qdrant.tech/>

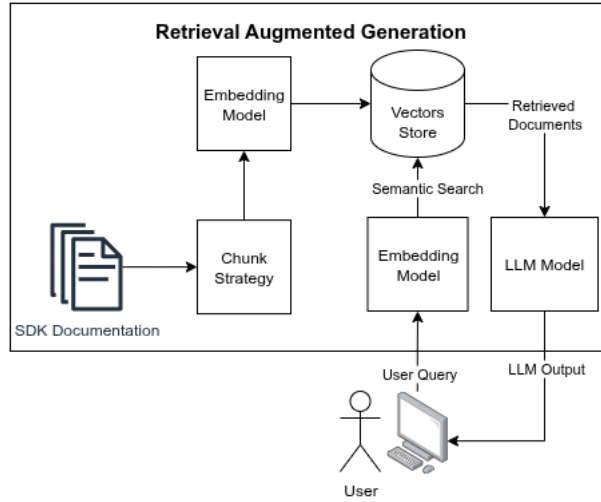


Fig. 2. Retrieval Augmented Generation - RAG

Once the prompt template is populated with the retrieved context, the LLM infers an answer.

To select an LLM, four generative models were compared: qwen2.5:7b [24], qwen3:8b [29], llama3.1:8b [30], and mistral:7b [31]. These models were chosen due to hardware restrictions and organizational policies for cloud applications. Open-source and smaller LLMs enabled on-premise deployment, ensuring private data security while balancing performance and computational efficiency.

3.3 Evaluation metrics

A combination of each chunk strategy with an embedding model and a generative LLM model has been implemented, resulting in 16 distinct experiments.

The Nomic embedding model was selected to generate vector embeddings for the LLM answer and ground truths for comparison within the RAGAS framework for semantic similarity. To compute BERTScore, the bertbase-uncased model was utilized.

The human evaluation consisted of 4 rounds, with only one round conducted per day for 12 evaluators, who were end users from the organization. The criteria for evaluating each generated answer were:

- Adequacy examines the quality of responses based on their completeness and richness;
- Usefulness judges the actual value of the answer helping to understand and perform the task effectively;
- Relevance measures the LLM response accuracy to the ground truth.

Each criterion was evaluated by developers using five-point Likert scale: Incorrect if the answer was totally incorrect; Inaccurate if the answer contained more wrong statements than correct ones; Acceptable if most statements were correct; Good if all statements were correct; Excellent if the answer was detailed and explained with examples beyond the ground truth.

4 Results and Discussion

Table 1 presents the evaluation metrics, with scores ranging from 0 to 1, where 0 represents the lowest and 1 the highest value. Based on these, we selected the best model.

LLama3.1:8B, Mistral:7B, and Qwen2.5:7B achieved their best scores using the split chunking type and the nomic-v1.5 embedding model, while Qwen3:8B performed best with the chunking sections type and bge-m3 embedding model.

Table 1. Overall scores

LLM	Embedding	Chunk	BLEU	ROUGE	Sem. Sim.	BERTScore
Mistral:7B	BGE-M3	Split	0.435	0.352	0.892	0.767
	BGE-M3	Section	0.442	0.347	0.884	0.765
	Nomic-V1.5	Split	0.435	0.353	0.888	0.767
	Nomic-V1.5	Section	0.461	0.349	0.888	0.767
Llama3.1:8B	BGE-M3	Split	0.409	0.357	0.886	0.774
	BGE-M3	Section	0.414	0.357	0.883	0.775
	Nomic-V1.5	Split	0.394	0.369	0.888	0.779
	Nomic-V1.5	Section	0.407	0.354	0.882	0.770
Qwen3:8B	BGE-M3	Split	0.411	0.399	0.899	0.778
	BGE-M3	Section	0.410	0.390	0.896	0.779
	Nomic-V1.5	Split	0.413	0.398	0.900	0.778
	Nomic-V1.5	Section	0.416	0.391	0.897	0.776
Qwen2.5:7B	BGE-M3	Split	0.470	0.383	0.897	0.784
	BGE-M3	Section	0.459	0.384	0.889	0.783
	Nomic-V1.5	Split	0.473	0.389	0.902	0.788
	Nomic-V1.5	Section	0.485	0.388	0.895	0.785

All models achieved higher BLEU scores using the section chunk method, while ROUGE and semantic similarity performed better with split chunks. Qwen3:8B model obtained its best score with section chunks for BERTScore, whereas the other three LLM models performed better with split chunks.

Among the embedding models, nomic-v1.5 achieved better results than BGE-M3. This indicates that, in the tested domain, the nomic-v1.5 vector of dimension 768 performed better than BGE-M3, which generated vectors of dimension 1024.

The best BLEU score of 0.485 was achieved by Qwen2.5:7B using the section chunk strategy and the nomic-v1.5 embedding model. Qwen3:8B, utilizing the bge-m3 embedding model and split chunk, attained the highest ROUGE score of 0.399. Qwen2.5:7B, with nomic-v1.5 and split chunk, secured the top score for semantic similarity 0.902 and BERTScore of 0.788.

To address RQ1, we tested various document chunking and embedding techniques. The combination of the split chunk method and the nomic-v1.5 embedding model slightly outperformed others, achieving the best scores in 3 out of 4 LLMs. This result shows that increasing the quantity of information per chunk does not always positively impact the LLM's response, as it may introduce irrelevant or redundant information, potentially causing hallucinations. Conversely, providing smaller pieces of information adds variety and enriches the LLM's context for improved responses.

For RQ2, we evaluated four different LLM models. The Qwen2.5:7B model achieved the highest scores in most metrics, except for the ROUGE score, where Qwen3:8B outperformed it with 0.39. However, Qwen2.5:7B excelled in BLEU (0.485), semantic similarity (0.902), and BERTScore (0.788).

Table 2 shows the metrics obtained by the Qwen2.5:7B model, paired with the nomic-v1.5 embedding model and split chunk method. Qwen2.5 achieved the highest scores in most metrics. The easy level obtained better scores than the medium level, and the model performed better for medium level questions than for hard level ones. The other three models exhibited similar behavior, with scores decreasing as questions became more chal-

lenging. This outcome was expected, as answering harder level questions required more refined and combined chunks, necessitating greater refinement to generate better answers.

Table 2. Model Scores by Question Level

LLM	Level	Embedding	Chunk	BLEU	ROUGE	Sem. Sim.	BERTScore
Mistral:7B	Easy	Nomic-V1.5	Split	0.443	0.364	0.895	0.770
	Medium	Nomic-V1.5	Split	0.427	0.340	0.886	0.766
	Hard	Nomic-V1.5	Split	0.410	0.313	0.850	0.751
Llama3.1:8B	Easy	Nomic-V1.5	Split	0.399	0.384	0.892	0.794
	Medium	Nomic-V1.5	Split	0.406	0.361	0.886	0.779
	Hard	Nomic-V1.5	Split	0.324	0.298	0.861	0.746
Qwen2.5:7B	Easy	Nomic-V1.5	Split	0.490	0.410	0.911	0.794
	Medium	Nomic-V1.5	Split	0.445	0.368	0.891	0.783
	Hard	Nomic-V1.5	Split	0.446	0.318	0.873	0.756
Qwen3:8B	Easy	BGE-M3	Section	0.425	0.411	0.905	0.789
	Medium	BGE-M3	Section	0.394	0.366	0.886	0.774
	Hard	BGE-M3	Section	0.351	0.315	0.865	0.734

The human evaluation average scores range from 1 to 5 and are presented in Table 3. This evaluation aimed to assess the answers generated by the optimal combination of chunk strategy and embedding model for each LLM model.

For human evaluation, we selected the combination that achieved the highest scores across most metrics for each LLM. Therefore, the selected models were: Llama3.1, Mistral, Qwen2.5 paired with split chunks and nomic-v1.5, and Qwen3 paired with section chunks and BGE-M3. The Qwen2.5:7B model achieved the highest average score across all three evaluated aspects, with the following average scores and their respective standard deviations (std.): Adequacy 3.27 (1.24), Relevance 3.36 (1.24), and Usefulness 3.36 (1.23).

For RQ3, Qwen2.5:7B also was the best model evaluated during the human evaluation phase, with a moderate standard deviation value, indicating that evaluators shared different opinions about generated answers. For overall scores, an adequacy score of 3.27 suggests that the model’s answer was sufficient to address the query. A relevance score of 3.36 indicates that the answer aligns with the defined ground truth. A usefulness score of 3.36 demonstrates that the model’s answer helps to solve the problem and its solutions are reasonable.

Table 4 shows Qwen2.5:7B human evaluation scores by question level. It is possible to notice that as the difficulty level increases, the scores decrease. This behavior is observed across all four models used in this experiment. This clarifies the tendency to reduce the score as questions become more difficult, a trend already mentioned in the metrics results.

Table 3. Human Evaluation - Overall Scores

Model	Adequacy	Relevance	Usefulness
Qwen2.5:7B	3.27	3.36	3.36
Llama3.1:8B	3.20	3.22	3.22
Mistral:7B	3.02	3.05	3.04
Qwen3:8B	2.97	2.99	2.99

Table 4. Human Evaluation - Model Scores by Question Level

LLM	Level	Adequacy	Relevance	Usefulness
Mistral:7B	Easy	3.17	3.19	3.18
	Medium	2.81	2.83	2.81
	Hard	2.71	2.72	2.74
Qwen2.5:7B	Easy	3.56	3.67	3.66
	Medium	2.82	2.85	2.90
	Hard	2.70	2.79	2.78
Qwen3:8B	Easy	3.24	3.26	3.26
	Medium	2.42	2.44	2.45
	Hard	2.81	2.85	2.86
Llama3.1:8B	Easy	3.37	3.37	3.37
	Medium	3.04	3.05	3.06
	Hard	2.64	2.68	2.68

5 Lessons Learned and Threats to Validity

During the human evaluation we did not measure the LLMs' hallucinations in generated answers due to the lack of volunteers to conduct this type of inspection manually. We were limited by computational resources to use an LLM as a judge.

To overcome computational limitations in evaluation, we adopted the BLEU and ROUGE metrics, which are widely used for NLP evaluations. However, they lack semantic understanding, as evidenced by their scores all below 0.5. In contrast, semantic similarity and BERTScore scores were all above 0.7. Therefore, adopting semantic similarity and BERTScore metrics was a reasonable choice to address the limitations of BLEU and ROUGE metrics.

The metrics values and human evaluation scores were very close to each other, indicating that the chunk strategies had no impact on embedding generation and, consequently, on the generated answers in the proposed scenarios.

6 Conclusion and Future Work

This paper investigated the application of the RAG technique in an organizational environment to assist software engineers to use a private SDK documentation for developing and maintaining automation scripts. Two distinct chunk methods were used to extract the documentation, which was then indexed in a vector database with two different embedding models, paired with four LLMs to generate answers to documentation-related questions. For each combination of chunk method, embedding model, and LLM, metrics were calculated, and human evaluation was conducted to assess aspects of the generated answers.

Given the limitations related to data privacy, organizational documentation, and computational resources, we chose to use on-premise small open-source LLMs to implement the RAG system. As shown in the results, a good perception of the generated answers was achieved, which is considered suitable for practical use and its integration with a larger system to provide this documentation assistance is also in course.

The split chunk method, combined with the nomic-v1.5 embedding model, performed better for three out of four LLMs. The only exception was Qwen3:8B, which performed better using BGE-M3 for ROUGE score and Qwen2.5:7B paired with section chunks achieved the highest score for BLEU metric. Qwen2.5:7B, when paired with split chunks and nomic-v1.5, achieved the highest scores for semantic similarity and BERTScore metrics. It also

achieved the best scores in human evaluations for the overall score across question levels, with scores tending to decrease as questions became more difficult. This necessitated improved information refinement in the indexing, retrieval, and generation phases, which will be addressed in future work.

For future work we intend to explore advanced RAG and modular RAG [6] techniques to improve results, the intention is that the system correctly answers most of hard level question. Additionally, tests with additional datasets to assess generalization for other documentations is also being considered.

Acknowledgment

This work is the result of the R&D project *Projeto de Engenharia de Software e Ciência de Dados aplicados ao Desenvolvimento de Sistemas*, performed by Sidia Instituto de Ciência e Tecnologia in partnership with Samsung Eletrônica da Amazônia Ltda., using resources from Federal Law No. 8.387/1991, and its disclosure and publicity are under the provisions of Article 39 of Decree No. 10.521/2020.

References

1. Kailash A Hambarde and Hugo Proenca. Information retrieval: recent advances and beyond. *IEEE Access*, 11:76581–76604, 2023.
2. Yutao Zhu, Huaying Yuan, Shuting Wang, Jiongnan Liu, Wenhan Liu, Chenlong Deng, Haonan Chen, Zheng Liu, Zhicheng Dou, and Ji-Rong Wen. Large language models for information retrieval: A survey. *arXiv preprint arXiv:2308.07107*, 2023.
3. Masoomali Fatehkia, Ji Kim Lucas, and Sanjay Chawla. T-rag: lessons from the llm trenches. *arXiv preprint arXiv:2402.07483*, 2024.
4. Nikhil Kandpal, Haikang Deng, Adam Roberts, Eric Wallace, and Colin Raffel. Large language models struggle to learn long-tail knowledge. In *International Conference on Machine Learning*, pages 15696–15707. PMLR, 2023.
5. Lei Huang, Weijiang Yu, Weitao Ma, Weihong Zhong, Zhangyin Feng, Haotian Wang, Qianglong Chen, Weihua Peng, Xiaocheng Feng, Bing Qin, et al. A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions. *ACM Transactions on Information Systems*, 43(2):1–55, 2025.
6. Yunfan Gao, Yun Xiong, Xinyu Gao, Kangxiang Jia, Jinliu Pan, Yuxi Bi, Yixin Dai, Jiawei Sun, Haofen Wang, and Haofen Wang. Retrieval-augmented generation for large language models: A survey. *arXiv preprint arXiv:2312.10997*, 2(1), 2023.
7. Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, et al. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in neural information processing systems*, 33:9459–9474, 2020.
8. Mahimai Raja, E Yuvarajan, et al. A rag-based medical assistant especially for infectious diseases. In *2024 International Conference on Inventive Computation Technologies (ICICT)*, pages 1128–1133. IEEE, 2024.
9. Muhammad Rafsan Kabir, Rafeed Mohammad Sultan, Fuad Rahman, Mohammad Ruhul Amin, Sifat Momen, Nabeel Mohammed, and Shafin Rahman. Legalrag: A hybrid rag system for multilingual legal information retrieval. *arXiv preprint arXiv:2504.16121*, 2025.
10. Md Saleh Ibtasham, Sarmad Bashir, Muhammad Abbas, Zulqarnain Haider, Mehrdad Saadatmand, and Antonio Cicchetti. Reqrag: Enhancing software release management through retrieval-augmented llms: An industrial study. In *International Working Conference on Requirements Engineering: Foundation for Software Quality*, pages 277–292. Springer, 2025.
11. Chengrui Wang, Qingqing Long, Meng Xiao, Xunxin Cai, Chengjun Wu, Zhen Meng, Xuezhi Wang, and Yuanchun Zhou. Biorag: A rag-llm framework for biological question reasoning. *arXiv preprint arXiv:2408.01107*, 2024.
12. Antonio Jimeno Yepes, Yao You, Jan Milczek, Sebastian Laverde, and Renyu Li. Financial report chunking for effective retrieval augmented generation. *arXiv preprint arXiv:2402.05131*, 2024.

13. Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, et al. A survey of large language models. *arXiv preprint arXiv:2303.18223*, 1(2), 2023.
14. Penghao Zhao, Hailin Zhang, Qinhan Yu, Zhengren Wang, Yunteng Geng, Fangcheng Fu, Ling Yang, Wentao Zhang, Jie Jiang, and Bin Cui. Retrieval-augmented generation for ai-generated content: A survey. *arXiv preprint arXiv:2402.19473*, 2024.
15. Ori Ram, Yoav Levine, Itay Dalmedigos, Dor Muhlgay, Amnon Shashua, Kevin Leyton-Brown, and Yoav Shoham. In-context retrieval-augmented language models. *Transactions of the Association for Computational Linguistics*, 11:1316–1331, 2023.
16. Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. Bleu: a method for automatic evaluation of machine translation. In *Proceedings of the 40th annual meeting of the Association for Computational Linguistics*, pages 311–318, 2002.
17. Chin-Yew Lin. Rouge: A package for automatic evaluation of summaries. In *Text summarization branches out*, pages 74–81, 2004.
18. Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q. Weinberger, and Yoav Artzi. Bertscore: Evaluating text generation with bert, 2020.
19. Tolga Şakar and Hakan Emekci. Maximizing rag efficiency: A comparative analysis of rag methods. *Natural Language Processing*, 31(1):1–25, 2025.
20. Irina Radeva, Ivan Popchev, Lyubka Doukovska, and Miroslava Dimitrova. Web application for retrieval-augmented generation: Implementation and testing. *Electronics*, 13(7):1361, 2024.
21. Yifan Yao, Jinhao Duan, Kaidi Xu, Yuanfang Cai, Zhibo Sun, and Yue Zhang. A survey on large language model (llm) security and privacy: The good, the bad, and the ugly. *High-Confidence Computing*, 4(2):100211, 2024.
22. Nir Kshetri. Cybercrime and privacy threats of large language models. *IT Professional*, 25(03):9–13, 2023.
23. Hanbo Huang, Yihan Li, Bowen Jiang, Lin Liu, Bo Jiang, Ruoyu Sun, Zhuotao Liu, and Shiyu Liang. On-premises llm deployment demands a middle path: Preserving privacy without sacrificing model confidentiality. In *ICLR 2025 Workshop on Building Trust in Language Models and Applications*, 2025.
24. An Yang, Baosong Yang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Zhou, Chengpeng Li, Chengyuan Li, Dayiheng Liu, Fei Huang, Guanting Dong, Haoran Wei, Huan Lin, Jialong Tang, Jialin Wang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Ma, Jianxin Yang, Jin Xu, Jingren Zhou, Jinze Bai, Jinzheng He, Junyang Lin, Kai Dang, Keming Lu, Keqin Chen, Kexin Yang, Mei Li, Mingfeng Xue, Na Ni, Pei Zhang, Peng Wang, Ru Peng, Rui Men, Ruize Gao, Runji Lin, Shijie Wang, Shuai Bai, Sinan Tan, Tianhang Zhu, Tianhao Li, Tianyu Liu, Wenbin Ge, Xiaodong Deng, Xiaohuan Zhou, Xingzhang Ren, Xinyu Zhang, Xipin Wei, Xuancheng Ren, Xuejing Liu, Yang Fan, Yang Yao, Yichang Zhang, Yu Wan, Yunfei Chu, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, Zhifang Guo, and Zhihao Fan. Qwen2 technical report, 2024.
25. Sriram Veturi, Saurabh Vaichal, Reshma Lal Jagadheesh, Nafis Irtiza Tripto, and Nian Yan. Rag based question-answering for contextual response prediction system. *arXiv preprint arXiv:2409.03708*, 2024.
26. Jianguo Wang, Xiaomeng Yi, Rentong Guo, Hai Jin, Peng Xu, Shengjun Li, Xiangyu Wang, Xiangzhou Guo, Chengming Li, Xiaohai Xu, et al. Milvus: A purpose-built vector data management system. In *Proceedings of the 2021 International Conference on Management of Data*, pages 2614–2627, 2021.
27. Vladimir Karpukhin, Barlas Öguz, Sewon Min, Patrick Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen tau Yih. Dense passage retrieval for open-domain question answering, 2020.
28. Zach Nussbaum, John X. Morris, Brandon Duderstadt, and Andriy Mulyar. Nomic embed: Training a reproducible long context text embedder, 2025.
29. An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao Ge, Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiayi Yang, Jing Zhou, Jingren Zhou, Junyang Lin, Kai Dang, Keqin Bao, Kexin Yang, Le Yu, Lianghao Deng, Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng Wang, Qin Zhu, Rui Men, Ruize Gao, Shixuan Liu, Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin, Xingzhang Ren, Xinyu Wang, Xinyu Zhang, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yinger Zhang, Yu Wan, Yuqiong Liu, Zekun Wang, Zeyu Cui, Zhenru Zhang, Zhipeng Zhou, and Zihan Qiu. Qwen3 technical report, 2025.
30. Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. The llama 3 herd of models. *arXiv e-prints*, pages arXiv-2407, 2024.
31. Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, Léo Re-

nard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas Wang, Timothée Lacroix, and William El Sayed. Mistral 7b, 2023.