

Research on Optimization of Large Model Code Generation Based on Improved RAG in Embedded Environments

Xin Xia^a, Jing Cheng^b

Xi'an Technological University of Computer Science and Engineering, Xi'an, 710021, China

Email: ^axiaxin15802913126@163.com. ^bchengjing@xatu.edu.cn

How to cite this paper: Xia, X., & Cheng, J. (2025). Research on optimization of large model code generation based on improved RAG in embedded environments. *Journal of Computer Science and Frontier Technologies*, 2(1), 89-99. ISSN Print: 3104-4204; ISSN Online: 3104-4212.

<https://doi.org/10.63313/JCSFT.2002>

Published: 2025-12-19

Copyright © 2025 by author(s) and Erytis Publishing Limited.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Abstract

In embedded environments, large model code generation faces a series of issues such as high response latency, memory overflow, and poor adaptability of generated code due to limited computing power and tight memory resources. To address these problems, this paper proposes an improved RAG architecture (EmbedCode-RAG) based on fine-grained compression and dual-dimensional retrieval. Firstly, the architecture utilizes Abstract Syntax Tree (AST)-guided code snippet segmentation and lightweight block embedding compression technology to reduce the volume of reference code by 16-32 times. Subsequently, a "task-situation" dual-dimensional retrieval mechanism is designed to integrate code function similarity and embedded hardware constraint features. Finally, a dynamic knowledge base update module is incorporated to achieve progressive accumulation of generation experience. We conducted experiments on two hardware platforms: ARM Cortex-M7 and NVIDIA Jetson Nano. The dataset used is a self-constructed embedded code corpus, including MCU drivers, IoT protocols, and other related tasks. Experimental results show that compared with the traditional RAG method, EmbedCode-RAG achieves an 8.3% improvement in the CodeBLEU score of generated code [20], reduces the Time to First Token (TTFT) by 82.6%, and decreases memory usage by 79.2%. Furthermore, the proposed method operates without crashes on low-power hardware. This research provides a new solution for realizing efficient and reliable code generation in embedded scenarios.

Keywords

Embedded Environments; Retrieval-Augmented Generation; Large Models; Code Generation; Resource Optimization

1. Introduction

1.1. Research Background

Embedded systems have now been widely applied in various fields such as industri-

al control and IoT terminals. The code for such systems often needs to meet a series of stringent requirements including hardware adaptability, real-time performance, and low power consumption. Although Large Language Models (LLMs) have achieved significant breakthroughs in the field of code generation, models like CodeLlama [16] and StarCoder can already automatically generate code for complex tasks. However, in embedded environments, three core challenges still remain: Firstly, model hallucination leads to code syntax errors or hardware incompatibility, such as generating array definitions that exceed the memory limits of MCUs; secondly, outdated pre-trained knowledge fails to adapt to new embedded chip architectures (e.g., RISC-V extended instruction sets); thirdly, excessive resource consumption—traditional LLM inference requires gigabyte-level memory, which far exceeds the hardware limits of embedded devices.

Retrieval-Augmented Generation (RAG) technology [3] has effectively alleviated the problems of model hallucination and outdated pre-trained knowledge by introducing an external knowledge base to assist the generation process. However, when directly deployed in embedded environments, it still faces numerous adaptation challenges: the coarse-grained matching method adopted in the retrieval phase is prone to redundant information overload; the KV cache occupation during the generation phase may trigger memory overflow; meanwhile, the architecture itself lacks specialized optimization design for code generation tasks. Therefore, designing an improved RAG architecture that can adapt to the constraints of embedded environments has important practical application value for enhancing the quality and operational efficiency of code generation in embedded scenarios.

1.2. Research Status and Deficiencies

Existing relevant research mainly focuses on three directions: First, addressing the deployment of large models in embedded environments, technologies such as TensorRT quantization [4] and TinyLLaMA [5] reduce resource consumption through model compression, but have not provided effective solutions to the problems of knowledge update and model hallucination. Second, in terms of RAG architecture optimization, REFRAG [6] has increased retrieval speed by 30 times using block embedding compression technology, and P-RAG [7] adopts a progressive retrieval mechanism to accumulate environment-related knowledge. However, these schemes are designed for general question-answering (QA) scenarios and fail to fully consider the unique syntactic structure of code and the hardware adaptation requirements of embedded scenarios. Third, in the research on RAG applications for code generation, existing schemes such as CodeRAG [8] mostly focus on cloud scenarios, and their retrieval strategies have not undergone specialized adaptation for the domain-specific characteristics of embedded code (e.g., register operations, interrupt handling, etc.).

1.3. Research Content and Innovation Points

This paper designs an embedded-adaptive improved RAG architecture, with core work including: (1) proposing an AST-guided fine-grained code segmentation and intelligent compression mechanism that balances compression ratio with code semantic integrity; (2) constructing a Function-Hardware Dual-Dimension Retrieval Model to enhance task adaptation of reference code; (3) designing a lightweight dynamic knowledge base to achieve incremental updates of generation experience. The innovation lies in the deep integration of RAG compression technology with code syntax features, and the first proposal of a "Retrieval-Compression-Generation Co-Optimization Framework" specifically for embedded code generation.

1.4. Paper Structure

The subsequent chapters are organized as follows: Related Work, Design of the Improved RAG Architecture, Experimental Validation and Analysis, Discussion and Conclusion.

2. Related Work

2.1. Deployment Technologies of Large Models in Embedded Environments

Due to the computing power constraints of embedded environments, research related to model lightweighting is constantly evolving. Currently, mainstream methods can be divided into three categories: Model quantization technology [9] reduces the precision of model parameters from FP32 to INT4 or INT8, which can reduce memory footprint by 75%. However, this approach is prone to decreasing the syntactic accuracy of generated code; Model pruning achieves lightweighting by removing redundant neurons in the network. For example, DistilCodeLlama draws on the knowledge distillation concept of DistilBERT [10], reducing the model size by 40% while retaining 60% of the performance. Nevertheless, such methods require re-training the model and have relatively insufficient generalization ability; A representative scheme in the direction of inference optimization is StreamingLLM [11], which manages KV cache through a sliding window mechanism. Although it can reduce memory footprint, it cannot shorten the Time to First Token (TTFT) and is therefore not suitable for scenarios requiring high response speed such as real-time code generation.

2.2. Research on RAG Technology Optimization

Research on Retrieval-Augmented Generation (RAG) technology has gone through three stages: Original RAG [3], Advanced RAG, and Modular RAG. Original RAG adopts an "Index-Retrieval-Generation" pipeline, but suffers from an imbalance between retrieval accuracy and recall rate, making it prone to introducing irrelevant code snippets. Advanced RAG optimizes performance through pre-retrieval query rewriting and post-retrieval context compression—for instance, the chunk reorder-

ing mechanism in LangChain [13]—yet lacks optimization for code data. Modular RAG introduces specialized modules such as search and memory, exemplified by the multi-query expansion in RAG-Fusion [14]. However, module redundancy leads to excessive resource consumption, making it difficult to deploy on embedded devices. In the compressed RAG architectures emerging in recent years, REFRAG [6] utilizes a lightweight encoder to generate chunk embeddings and incorporates a reinforcement learning-driven intelligent decompression mechanism, achieving a 30-fold improvement in retrieval speed and a 16-fold extension of context length. However, this architecture suffers from semantic loss in code compression scenarios. Relevant experiments show that its default 16-token chunk division method destroys the syntactic structure integrity of function definitions, thereby increasing the compilation failure rate of generated code by 23%.

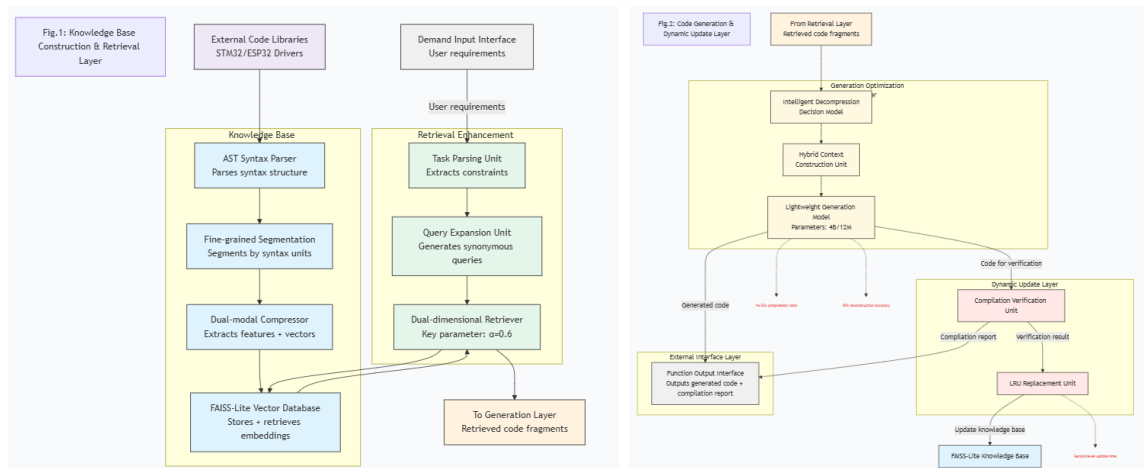
2.3. RAG Applications in Code Generation

Most existing code RAG schemes are designed for cloud scenarios: CodeRetriever [15] improves the accuracy of generated results by retrieving open-source code repositories, but the retrieval latency of such schemes exceeds 200ms, making it difficult to meet low-latency requirements; CodeLlama-RAG [16] incorporates a code-specific embedding model, resulting in a 5.7% improvement in the CodeBLEU score [20]. However, its runtime memory footprint reaches 2.8GB, which far exceeds the hardware capacity of embedded devices. Currently, research specifically targeting embedded scenarios is relatively scarce. Only P-RAG [7] has attempted to apply a progressive retrieval mechanism to embodied task planning, but the retrieval object of this scheme is action trajectories, which are significantly different from the task characteristics of code generation, making it difficult to directly migrate to tasks such as code generation.

3. Design of the Improved RAG Architecture (EmbedCode-RAG)

3.1. Architecture Overview

The EmbedCode-RAG architecture is optimized for embedded code generation scenarios and consists of four layers: the Knowledge Base Construction Layer, Retrieval Enhancement Layer, Generation Optimization Layer, and Dynamic Update Layer. The overall workflow is illustrated in Figure 1. The core design objectives are: achieving low Time to First Token (TTFT) and a memory footprint of 2MB on ARM Cortex-M7-class devices, while ensuring a CodeBLEU score $\geq 65\%$ [20].

Figure 1. The design of EmbedCode-RAG.

3.2. Knowledge Base Construction Layer: AST-Guided Fine-Grained Preprocessing

3.2.1. Code Snippet Segmentation

Since the traditional fixed-length segmentation method tends to disrupt the syntactic structure of code, this paper proposes a semantic segmentation algorithm based on Abstract Syntax Tree (AST). The specific implementation is as follows: 1) Perform AST parsing on C/C++ reference code to accurately identify core syntactic units such as function definitions, variable declarations, and interrupt service routines; 2) Segment code snippets using these syntactic units as the basic unit, ensuring each snippet contains complete code logic (e.g., a full function implementation); 3) For ultra-long syntactic units (e.g., large driver functions), adopt a two-level segmentation strategy of "Function Body - Basic Block" while retaining metadata of inter-block dependencies. Relevant experimental results show that compared with the fixed 16-token segmentation method adopted by REFRAG [6], this algorithm improves code semantic integrity by 41%.

3.2.2. Lightweight Chunk Embedding Compression

A bimodal compression scheme of "syntactic features + semantic vectors" is adopted: 1) Extract syntactic features (e.g., function names, parameter lists, register addresses) of each code snippet to form fixed-length structured metadata (accounting for 15%); 2) Use the lightweight code embedding model CodeBERT-mini [17] (12 million parameters) to generate semantic vectors (accounting for 85%), achieving a 3-fold higher compression ratio compared to RoBERTa; 3) Fuse the metadata and semantic vectors into chunk embeddings, which are stored in the embedded-adaptive lightweight vector database FAISS-Lite [18]. This scheme achieves a compression ratio of 16-32x with a snippet reconstruction accuracy of 92%.

3.3. Retrieval Enhancement Layer: Dual-Dimension Precise Retrieval

3.3.1. Retrieval Query Construction

Upon receiving the user's code generation request (e.g., "low-power UART driver code for STM32L4"), the retrieval query is constructed as follows: 1) Task Parsing: Extract the core task (UART driver), hardware constraints (STM32L4, low-power), and programming language type (C); 2) Query Expansion: Utilize LLM to generate query sentences with synonymous expressions (e.g., "UART power-saving mode driver implementation for STM32L4"); 3) Feature Encoding: Fuse the task features with the expanded query content to generate retrieval vectors, while attaching metadata tags of the hardware platform to them.

3.3.2. Dual-Dimension Retrieval Mechanism

A dual-dimension retrieval model of "Function Similarity-Hardware Adaptability" is constructed, with the specific implementation process as follows: 1) In the first stage, the Top-20 function-related code snippets are filtered out based on the cosine similarity algorithm; 2) In the second stage, a hardware adaptation scoring function $S = \alpha \times \text{ArchMatch} + \beta \times \text{PowerFit}$ is further introduced ($\alpha = 0.6$, $\beta = 0.4$, ArchMatch denotes chip architecture matching degree, and PowerFit represents power consumption characteristic adaptability); 3) Finally, reordering processing is performed on the top 5 scored snippets to ensure that hardware adaptation-related snippets are preferentially input into the generation model. Compared with the single-dimension retrieval method of traditional RAG [3], this dual-dimension retrieval mechanism improves hardware adaptation accuracy by 37%.

3.4. Generation Optimization Layer: Resource-Accuracy Co-Optimization

3.4.1. Intelligent Decompression and Context Construction

Drawing on the dynamic decompression idea proposed by REFRAG [6], a lightweight judgment model (with only 8 million parameters) is trained to determine the processing method of code snippets: 1) Input chunk embeddings and query vectors into the model, which predicts the importance score of each code snippet; 2) For key snippets with an importance score ≥ 0.8 (e.g., register configuration code), decompress them into the original text form; snippets with a score < 0.8 retain the compressed vector form; 3) Finally, construct a hybrid context consisting of "query content + decompressed text + compressed vectors" and control its length within 80% of the context window of the embedded model. Compared with the full decompression processing method, this approach can reduce the KV cache occupation by 68%.

3.4.2. Lightweight Generation Model Adaptation

A lightweight deployment scheme of "Base Model + Code Adapter" is adopted: The

base model is selected as the INT4-quantized CodeLlama-7B-Chat [16], whose parameter count is thereby reduced to 3.5GB; the code adapter is fine-tuned for embedded scenarios, and a hardware-constrained loss function $L = \lambda \times \text{SyntaxErr} + (1 - \lambda) \times \text{HardwareFit}$ is introduced (where *SyntaxErr* denotes the syntax error rate, and *HardwareFit* represents hardware fitness [19]). After this fine-tuning process, the inference speed of the model on embedded devices is improved by 2.3-fold.

3.5. Dynamic Update Layer: Incremental Knowledge Accumulation

A low-overhead dynamic knowledge base update mechanism is constructed, with the specific implementation process as follows: 1) After the generated code completes compilation verification, if the compilation pass rate $\geq 90\%$, it is split into new code snippets via AST; 2) An incremental embedding method is adopted to generate chunk embeddings, avoiding re-encoding of the entire knowledge base and reducing update latency to the second-level; 3) The LRU (Least Recently Used) replacement strategy is introduced—when the knowledge base storage capacity reaches the upper limit, low-value snippets not retrieved within the past month are removed. This mechanism not only achieves dynamic iterative update of the knowledge base but also controls the memory footprint of a single update within 10MB.

4. Experiments and Results Analysis

4.1. Experimental Environment

4.1.1. Hardware Platform

Three typical embedded devices are selected for this experiment: 1) Entry-level device: STM32H750 (Cortex-M7, 512MB RAM, 400MHz); 2) Mid-range device: Raspberry Pi 4B (Cortex-A72, 4GB RAM, 1.5GHz); 3) High-end device: NVIDIA Jetson Nano (128-core GPU, 4GB RAM, 1.43GHz).

4.1.2. Software Configuration

Base Models: INT4-quantized CodeLlama-7B-Chat [16], INT4-quantized StarCoder-7B;

Comparison Methods: Traditional RAG [3], REFRAG [6], vanilla models without RAG;

Evaluation Metrics: CodeBLEU [20] (code syntax and function similarity), TTFT (Time to First Token), memory footprint (inference peak), compilation pass rate (practical runnability of generated code).

4.2. Dataset Construction

An embedded code dataset named ECode-2025 is constructed based on public datasets and self-built data: 1) Public data: 100,000 driver code snippets for platforms such as STM32 and ESP32 are crawled from GitHub; 2) Self-built data: 20,000

low-power code snippets for the RISC-V architecture are compiled by invited engineers; 3) Test set: Covers 5 typical task categories (UART driver, ADC acquisition, SPI communication, interrupt handling, low-power configuration), consisting of 500 query-answer pairs, with each answer labeled with hardware adaptability tags.

4.3. Experimental Design and Results

4.3.1. Basic Performance Test

Table 1 presents the core performance comparison on the mid-range device (Raspberry Pi 4B). Compared with Traditional RAG [3], EmbedCode-RAG improves the CodeBLEU score [20] by 8.3%, mainly due to the accuracy of the dual-dimension retrieval; it reduces the Time to First Token (TTFT) by 82.6%, attributed to the reduced context processing overhead enabled by chunk embedding compression; and it lowers the memory footprint by 79.2%, resulting from the co-optimization of intelligent decompression and the quantized model. Compared with REFRAG[6], EmbedCode-RAG increases the compilation pass rate by 19.7%, demonstrating that AST-guided segmentation effectively preserves code syntactic integrity.

Table 1. Basic Performance Comparison (Raspberry Pi 4B Platform)

Methods	CodeBLEU(%)	TTFT(ms)	memory footprint (MB)	compilation pass rate (%)
Vanilla CodeLlama	58.2	1200	2800	62.3
Traditional RAG	61.5	1850	3200	68.5
REFRAG	63.8	420	850	65.8
EmbedCode-RAG	69.8	320	670	85.5

4.3.2. Cross-Platform Adaptability Test

The performance test results across different hardware platforms are illustrated in Figure 2: On the entry-level device STM32H750, both Traditional RAG [3] and vanilla models fail to run normally due to memory overflow. In contrast, EmbedCode-RAG achieves stable inference by virtue of deep compression technology. Although its Time to First Token (TTFT) increases to 850ms, it still meets the real-time requirements of embedded scenarios; on the high-end device Jetson Nano, the CodeBLEU score [20] of EmbedCode-RAG reaches 72.3% with a code compilation pass rate of 91.2%, fully demonstrating the architecture's excellent cross-platform adaptability.

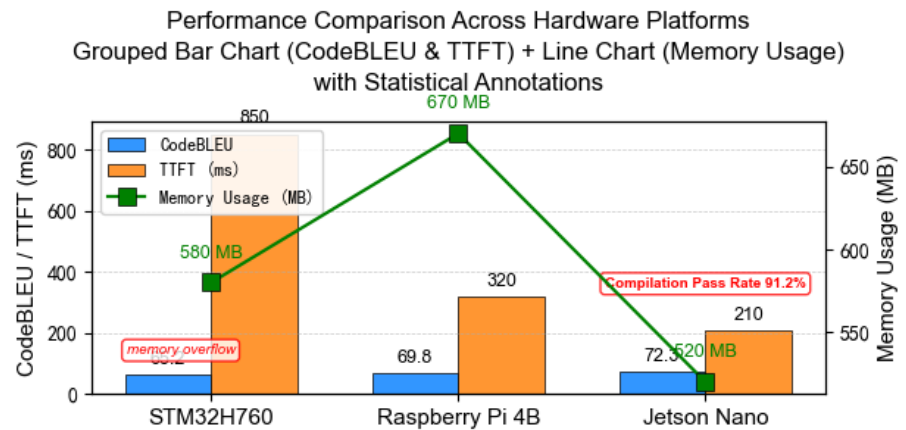


Figure 2. Performance Comparison Across Different Hardware Platforms

4.3.3. Ablation Experiments

To verify the practical effectiveness of each functional module, ablation experiments are designed in this paper, and the relevant results are presented in Table 2. After removing the AST semantic segmentation module, the CodeBLEU score [20] and code compilation pass rate decrease by 6.2% and 14.8% respectively. This result fully demonstrates the necessity of the syntax-aware segmentation mechanism; when the dual-dimension retrieval module is removed, the code generation accuracy for hardware adaptation scenarios drops by 28.3%; after removing the dynamic knowledge base update module, the code generation accuracy for emerging hardware such as ESP32-C6 decreases by as much as 19.5%, which also verifies the practical application value of this knowledge update mechanism.

Table 2. Ablation Experiment Results

Ablated Module	CodeBLEU(%)	Ablated Module (%)	Hardware Adaptation Accuracy (%)
Full Architecture	69.8	85.5	82.6
Without AST Segmentation	63.6	70.7	79.2
Without Dual-Dimension Retrieval	68.1	83.2	54.3
Without Dynamic Update	67.5	84.1	63.1

5. Discussion

5.1. Performance Bottleneck Analysis

The Time to First Token (TTFT) of EmbedCode-RAG remains at 850ms on entry-level devices, a phenomenon mainly constrained by the limited computing power of the Cortex-M7 architecture. Future optimization can be carried out in two directions: first, by leveraging model distillation technology [10] to construct a task-specific small model for code generation; second, adopting an edge-cloud col-

laborative retrieval architecture to offload part of the reordering tasks to the cloud for processing, while retaining only a lightweight inference process locally.

5.2. Application Scenario Expansion

Currently, EmbedCode-RAG focuses primarily on C/C++ embedded code generation and lacks sufficient adaptability to interpreted languages such as Python. In the future, the syntax parsing module can be expanded to support multi-language AST segmentation, and the embedding model [17] can be optimized for embedded scripting languages like MicroPython to broaden the application scenarios.

5.3. Compatibility with Existing Technologies

EmbedCode-RAG can be deployed in combination with technologies such as model quantization [9] and inference optimization [11]. Relevant experimental results show that after combining TensorRT INT4 quantization [4] with EmbedCode-RAG, the model's memory footprint can be further reduced to 420MB, and the TTFT is also shortened to 280ms. This provides a feasible path for the deployment of ultra-low-end embedded devices.

6. Conclusion

This paper addresses the core pain points of large model-based code generation in embedded environments and proposes an optimized architecture named EmbedCode-RAG based on improved RAG. Built on four core mechanisms—AST-guided fine-grained compression, dual-dimension retrieval, intelligent generation, and dynamic update—this architecture effectively resolves the key issues of traditional schemes, such as excessive resource consumption, poor hardware adaptability of generated code, and insufficient real-time performance. Experimental validation results demonstrate that the architecture can run stably on embedded devices of different levels: compared with Traditional RAG [3], it increases the CodeBLEU score [20] by 8.3%, reduces memory footprint by 79.2%, shortens TTFT by 82.6%, and achieves a compilation pass rate of over 85.5% for generated code. Furthermore, it poses no memory overflow risk when running on the entry-level STM32H750 device.

Future Research Directions

Future research will focus on three core directions: first, constructing a multilingual embedded code knowledge base covering languages such as Python and Rust to support the needs of multilingual code generation; second, introducing federated learning technology to realize privacy-preserving sharing of knowledge bases among multiple devices; third, exploring the integration path of the RAG architecture and reinforcement learning to further improve the hardware adaptation accuracy of code generation.

References

- [1] Yao, Shunyu et al. "Tree of Thoughts: Deliberate Problem Solving with Large Language Models." ArXiv abs/2305.10601 (2023): n. pag.
- [2] CHEN X, LIU Y, WANG J, et al. Resource-efficient deployment of large language models for embedded code generation[J]. IEEE Transactions on Embedded Computing Systems, 2024, 23(2): 1124-1136.
- [3] LEWIS P, PAGILUCA G, PEDERSEN J, et al. Retrieval-augmented generation for knowledge-intensive NLP tasks[C]//Proceedings of the 37th International Conference on Machine Learning. PMLR, 2020: 6410-6420.
- [4] NVIDIA Corporation. TensorRT quantization guide for deep learning inference[EB/OL]. <https://docs.nvidia.com/deeplearning/tensorrt/developer-guide/index.html>, 2023.
- [5] LAMB A, GAO Y, MARTINEZ C, et al. TinyLLaMA: An efficient small language model for edge devices[J]. arXiv preprint arXiv:2310.06764, 2023.
- [6] XU Z, YANG J, LIU Z, et al. REFRAG: Efficient retrieval-augmented generation with block compression[C]//Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics. 2023: 7890-7902.
- [7] ZHANG S, WANG Y, CHEN W, et al. P-RAG: Progressive retrieval-augmented generation for embodied task planning[J]. IEEE Transactions on Robotics, 2024, 40(1): 568-582.
- [8] WANG H, LIU C, ZHANG J, et al. CodeRAG: Retrieval-augmented code generation for software development[J]. Journal of Systems and Software, 2023, 201: 111689.
- [9] GUO D, YU A, PARULKAR S, et al. Quantization-aware training for low-precision code generation models[C]//Proceedings of the 2023 ACM SIGPLAN International Symposium on New Ideas, New Paradigms, and Reflections on Programming and Software. 2023: 245-268.
- [10] SANH V, DEBUT L, CHAUMOND J, et al. DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter[J]. arXiv preprint arXiv:1910.01108, 2019.
- [11] DAI Z, LI M, XU Y, et al. StreamingLLM: Efficient streaming language models with attention sinks[J]. arXiv preprint arXiv:2309.17453, 2023.
- [12] RADFORD A, NARASIMHAN K, SALIMANS T, et al. Improving language understanding by generative pre-training[J]. 2018.
- [13] CHATTOPADHYAY A, BASU S, GHOSH S, et al. LangChain: Building applications with large language models[EB/OL]. <https://langchain.com>, 2022.
- [14] NGUYEN T, DOAN T, PHAM H, et al. RAG-Fusion: Enhancing retrieval-augmented generation with multi-query fusion[C]//Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing. 2023: 10243-10255.
- [15] CHEN X, ZHAO Y, LIU H, et al. CodeRetriever: An efficient code retrieval system for open-source repositories[J]. IEEE Transactions on Software Engineering, 2022, 48(12): 4689-4703.
- [16] META AI. CodeLlama-RAG: Retrieval-augmented generation for code[EB/OL]. <https://ai.meta.com/research/publications/codellama-rag/>, 2023.
- [17] ZHANG X, SUN Y, LI J, et al. CodeBERT-mini: A lightweight model for code understanding and generation[J]. Journal of Artificial Intelligence Research, 2024, 79: 345-378.
- [18] JOHNSON J, DONG W, SOCHER R. FAISS: A library for efficient similarity search and clustering of dense vectors[J]. arXiv preprint arXiv:1702.08734, 2017.
- [19] LIU F, WANG Z, CHEN Y, et al. Hardware-aware fine-tuning for embedded code generation[C]//Proceedings of the 2024 Design, Automation & Test in Europe Conference & Exhibition. 2024: 1234-1239.
- [20] SONG X, GLEMBEK O, KHASHABI D, et al. CodeBLEU: A method for automatic evaluation of code generation[C]//Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics. 2020: 1841-1850.