



ISSUE  
Nº94



EUROPEAN OPEN  
SCIENCE SPACE

COLLECTION OF SCIENTIFIC PAPERS



8<sup>TH</sup> INTERNATIONAL  
SCIENTIFIC  
AND PRACTICAL  
CONFERENCE

EVOLVING SCIENCE:  
THEORIES, DISCOVERIES  
AND PRACTICAL  
OUTCOMES

JUNE 29 – JULY 1, 2026, ZURICH, SWITZERLAND



# AI-DRIVEN INFORMATION SYSTEMS DEVELOPMENT: LLM-BASED CODE GENERATION APPROACHES

Знахур Людмила

ст.викладач

Кафедра інформаційних систем

Харківський національний економічний університет

імені Семена Кузнеця, Україна

## Abstract

The automation of software code generation has emerged as one of the most consequential applications of large language models (LLMs), fundamentally altering the productivity landscape of information system engineering. This article presents a comprehensive review and empirical analysis of LLM-based code generation systems, covering the period from 2022 to 2026, and proposes a hybrid architecture that combines Retrieval-Augmented Generation (RAG) with a Mixture-of-Experts (MoE) framework to address persistent limitations in existing approaches. Following systematic review principles, the study surveys model architectures — including left-to-right autoregressive decoders, masked language models, and encoder-decoder designs — alongside leading benchmarks such as HumanEval, SWE-Bench, MBPP, and LiveCodeBench. Code quality analysis using static analysis tools shows that the hybrid model achieves higher readability and lower cyclomatic complexity than both baselines. The article identifies critical gaps, including hallucination rates, security vulnerabilities in generated code, and limited support for low-resource programming languages, and proposes targeted improvement strategies involving self-reflection loops, post-generation verification, and domain-adaptive fine-tuning. The findings underscore the necessity of multi-component evaluation frameworks for production-grade LLM code generation deployments.

**Keywords:** large language models, code generation, retrieval-augmented generation, mixture of experts, software engineering automation, HumanEval, SWE-Bench, transformer architectures, automated programming, benchmark evaluation.

## Introduction

Software engineering has long pursued the goal of automating the labor-intensive and error-prone process of writing code. Early program synthesis efforts, dating to the 1960s, employed rule-based and template-driven systems that produced code from formal specifications, but these approaches were rigid and could not generalise beyond narrowly defined task domains. The emergence of machine learning, and subsequently deep learning, transformed this trajectory by enabling models to learn statistical patterns from vast repositories of human-authored code, progressively closing the gap between specification and implementation.

The most significant inflection point occurred with the introduction of large language models pre-trained on both natural language and source code. Models such as OpenAI's Codex and the derivative GitHub Copilot demonstrated for the first time that a single neural model could translate natural-language descriptions into functional

code across multiple programming languages, achieving performance levels previously considered unattainable by automated systems. This catalysed an extraordinary proliferation of research and commercial activity: in 2025, 84% of professional developers routinely employ AI-powered coding assistants, and the market for AI code generation software is projected to grow at a compound annual rate exceeding 20%, reaching multi-billion valuations by 2030.

Despite these advances, substantial challenges remain unresolved. LLM-generated code is subject to hallucinations — the generation of syntactically plausible but semantically incorrect or functionally broken code — a phenomenon that 46% of developers cite as a primary concern in adoption surveys. Security vulnerabilities in generated code represent a related and more severe problem: benchmark evaluations on SEC-Bench reveal persistent vulnerability rates in the outputs of even the most capable models. Generalization failures across programming languages, particularly those underrepresented in training corpora, further limit the practical applicability of current systems. Finally, the computational and financial costs of frontier commercial models restrict accessibility for small enterprises and individual developers.

This article addresses these challenges through two complementary contributions. First, it provides a systematic synthesis of existing model architectures, evaluation benchmarks, and performance results for LLM-based code generation systems as of 2025, establishing an evidence-based baseline for understanding the current state of the field. Second, it proposes and empirically evaluates a hybrid code-generation architecture that integrates RAG with an MoE routing strategy, supplemented by iterative self-reflection, and demonstrates measurable improvements in functional correctness and code quality relative to established baselines.

### **Aim and objectives of the research**

The primary aim of this research is to investigate the capabilities and limitations of large language model-based code generation systems and to develop a theoretically grounded and empirically validated hybrid architecture that improves functional correctness, code quality, and security relative to existing baseline approaches. The research pursues the following specific objectives:

1. To systematically analyze the architectural landscape of LLM-based code generation, covering autoregressive, masked, and encoder-decoder paradigms, and to situate each within the broader trajectory of model development from 2022 to 2026.
2. To conduct a structured comparative evaluation of leading code generation models — including Claude 3.7 Sonnet, GPT-4o, Gemini 2.5 Pro, DeepSeek V3, and Llama.
3. To design a hybrid architecture combining RAG for context-aware code retrieval and MoE for task-specialized computation, incorporating self-reflection mechanisms for iterative output refinement.
4. To implement a Python-based prototype of the proposed architecture using the QZhou-Embedding model, the FAISS vector database, and transformer-based generation components, and to evaluate the prototype on the HumanEval benchmark.

5. To assess code quality using quantitative static analysis metrics — pylint scoring, cyclomatic complexity, and the Maintainability Index — and to compare the hybrid model's outputs against those of GPT-4o and CodeLlama.

### **Literature review**

The history of automated code generation spans six decades, during which the field underwent a series of discrete paradigm shifts driven by advances in computational theory, formal methods, and machine learning. The earliest programme synthesis systems of the 1960s and 1970s relied on formal specifications and deductive reasoning to derive code from mathematical descriptions of intended behaviour. These systems were theoretically rigorous but practically limited: they required specifications that were as precise and complete as the code itself, effectively displacing rather than reducing the intellectual burden on developers.

The decisive transition occurred between 2020 and 2022 with the publication of GPT-3 and Codex, which demonstrated that transformer architectures pre-trained on internet-scale text could be fine-tuned or prompted to generate syntactically correct and contextually appropriate code at a quality level that matched or exceeded domain-specialist models [1-7]. This established the large language model paradigm as the dominant framework for code generation research, a position it has maintained and reinforced through subsequent generations of increasingly capable systems. Contemporary code generation models are distributed across three principal architectural families, each associated with characteristic strengths and limitations relative to specific downstream tasks. Left-to-right autoregressive models, which predict each token conditioned on all preceding tokens, constitute the dominant paradigm for code generation tasks that require the sequential construction of complete programmes. Models such as CodeGPT, CodeParrot, Codex, GPT-NeoX, and Google's PaLM-Code exemplify this family. Their unidirectional generation process maps naturally onto the common developer workflow of writing code from top to bottom, and the task of predicting the next token is computationally well-suited to parallel training on large corpora. The principal limitation of this family is the difficulty of incorporating information that appears later in the context — so-called right-context — into the generation of earlier tokens, which is especially relevant when code structure is determined by patterns only revealed as the programme unfolds. Masked language models, which predict randomly masked tokens from bidirectional context, are better suited to representation learning and classification tasks than to generation. CodeBERT and CuBERT are the canonical examples of this family in the code domain. Because their bidirectional attention mechanism enables each token representation to incorporate context from both preceding and following tokens, they produce richer semantic embeddings than autoregressive models of comparable scale. These embeddings support downstream tasks such as code clone detection, defect prediction, and code search, but the models are not directly applicable to left-to-right code generation without architectural modification. Encoder-decoder models combine a bidirectional encoder with an autoregressive decoder, enabling the model to construct a rich representation of an input sequence and then generate an output sequence conditioned on that representation. Pre-training objectives such as masked span prediction and denoising for sequence reconstruction make these models well-suited to conditional generation tasks

of natural language-to-code, code summarisation, code translation across programming languages, and comment generation. CodeT5 and PLBART are the primary representatives of this family, achieving strong performance on conditional generation benchmarks while remaining more parameter-efficient than comparable purely autoregressive models (Table 1).

Table 1. Evolution of Key Model Types in Code Generation (2022–2026)

| Year | Model Type                            | Representative Models                                   | Key Innovations                                                                                                                                              |
|------|---------------------------------------|---------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 2022 | Early LLMs                            | Codex, AlphaCode                                        | Sequence-to-sequence generation; initial benchmark success on HumanEval [7]                                                                                  |
| 2023 | Code-Focused LLMs                     | Code LLaMA, StarCoder                                   | Open-source training on large code corpora; improved syntax handling                                                                                         |
| 2024 | Agentic Systems                       | Devin, ChatDev                                          | Multi-agent collaboration; tool integration; autonomous task decomposition [2,4, 7]                                                                          |
| 2025 | Advanced Hybrids                      | Claude 3.7, Gemini 2.5                                  | Ultra-long context windows; self-refinement loops; security-aware generation [8, 9]                                                                          |
| 2026 | Frontier Agentic & Specialized Models | Claude Opus 4.7, GPT-5.5, Kimi K2.5, GLM-5, Qwen3-Coder | Formal verification integration; Terminal-Bench as a standard for real-world agentic workflows; open-weight models rivalling commercial frontier systems [9] |

The maturation of code generation research has been accompanied by the development of increasingly sophisticated evaluation frameworks that go beyond surface-level syntactic similarity measures to assess functional correctness, real-world applicability, and code quality. The HumanEval dataset established the pass@k metric as the primary measure of functional correctness for code generation: a response is considered correct if at least one of the k generated candidates passes all provided unit tests. HumanEval's 164 Python programming problems span a range of algorithmic and data manipulation tasks and remain the most widely reported benchmark for function-level code generation.

BLEU and CodeBLEU metrics, adapted from machine translation evaluation, measure n-gram overlap between generated and reference code. While computationally convenient, these metrics have been shown to correlate poorly with functional correctness and developer judgments of code quality, as semantically equivalent implementations may share few surface-level tokens. The CodeBLEU extension incorporates abstract syntax tree matching and data flow analysis to capture more structural similarity, but remains an imperfect proxy for execution-based correctness.

SWE-Bench, introduced in 2024, represents a qualitative advance in benchmark design by moving from isolated function-level tasks to full software engineering problems extracted from real GitHub issue trackers. Resolving a SWE-Bench instance requires a model to understand an existing codebase, identify the location of a defect

or missing feature, generate a corrective patch, and produce output that passes the repository's automated test suite. Resolution rates on SWE-Bench are substantially lower than HumanEval pass rates, exposing a significant gap between laboratory performance and real-world applicability.

The benchmark landscape will expand further to include LiveCodeBench for dynamic competitive programming assessment, SEC-Bench for security-aware vulnerability detection, and RepoBench for repository-level dependency understanding (Table 2). Evaluation metrics correspondingly now encompass computational efficiency (API call counts, latency), code maintainability (cyclomatic complexity, Maintainability Index), and security (hallucination rates, vulnerability counts).

Table 2. Typical Performance Ranges by Code Generation Task Category

| Task Category                 | Typical Pass@k / Accuracy Range | Primary Challenges                                                |
|-------------------------------|---------------------------------|-------------------------------------------------------------------|
| Code Completion               | 80–95%                          | Context understanding; long-range dependency resolution           |
| NL-to-Code (simple functions) | 60–85%                          | Ambiguity in natural language; translation of complex logic       |
| Program Synthesis             | 40–70%                          | Formal specification understanding; ensuring logical correctness  |
| Code Translation              | 70–90%                          | Semantic preservation; idiomatic translation across languages     |
| Automated Bug Fixing          | 30–60%                          | Root cause identification; generating minimal and correct patches |

The leading code generation models of 2025-2026 represent a diverse ecosystem in which commercial frontier models, open-source alternatives, and hybrid agentic systems coexist and complement one another (Table 3-5). Claude 3.7 Sonnet, developed by Anthropic, achieves the highest SWE-Bench resolution rate among evaluated models at 70.3%, attributable to its superior multi-file context handling and reasoning capabilities within a 200,000-token context window. GPT-4o, developed by OpenAI, offers a balance of speed, multimodal capability, and ecosystem integration that supports broad developer adoption, but exhibits variable performance on repository-level tasks. Gemini 2.5 Pro, with a context window of up to two million tokens, excels at dependency analysis and whole-repository reasoning tasks, representing a different capability profile optimized for scale rather than single-query speed. Open-source models, including DeepSeek V3, Llama 4 Maverick, and Qwen 2.5 Coder have substantially narrowed the performance gap with commercial models on function-level benchmarks while offering cost and customization advantages. DeepSeek V3 employs a Mixture of Experts architecture that achieves computational efficiency by routing each token to a small subset of specialized expert modules, enabling strong HumanEval performance at a fraction of the inference cost of dense-parameter commercial models. Llama 4 Maverick, with context windows of up to

ten million tokens, is designed for edge deployment and creative coding tasks, though it underperforms commercial models on complex reasoning tasks.

Table 3. Benchmark Performance Comparison of Leading Models (2025)

| Benchmark              | Claude 3.7 Sonnet | GPT-4o | Gemini 2.5 Pro | DeepSeek V3 |
|------------------------|-------------------|--------|----------------|-------------|
| HumanEval (Pass@1)     | 92%               | 90%    | 89%            | 90%         |
| SWE-Bench (% Resolved) | 70%               | 33–69% | ~64%           | 49%         |
| LiveCodeBench (Pass@1) | ~50%              | 30–79% | ~70%           | ~64%        |

Table 4. HumanEval Pass@1 Scores for Leading Code Generation Models (2025)

| Model             | Pass@1 (%) | Notable Characteristics                                    |
|-------------------|------------|------------------------------------------------------------|
| Claude 3.7 Sonnet | 92         | High reasoning capability; low hallucination rate          |
| GPT-4o            | 90         | Multimodal inputs; fast inference pipeline                 |
| Gemini 2.5 Pro    | 89         | Massive context window; strong structured-task performance |
| DeepSeek V3       | 90         | Cost-effective Mixture of Experts architecture             |
| Llama 4 Maverick  | 62         | Open-source; suited to creative and edge-device tasks      |

Table 5. Comprehensive Model Comparison Across Multiple Benchmarks (2025)

| Model             | HumanEval (Pass@1) | MBPP (Acc.) | SWE-Bench (%) | LiveCodeBench | Cost (\$/M tokens) |
|-------------------|--------------------|-------------|---------------|---------------|--------------------|
| Claude 3.7 Sonnet | 92%                | ~85%        | 70.3%         | ~50%          | \$15               |
| GPT-4o            | 90%                | ~90%        | 33–69%        | 30–79%        | \$10               |
| Gemini 2.5 Pro    | 89%                | ~88%        | ~64%          | ~70%          | \$12               |
| DeepSeek V3       | 90%                | N/A         | 49%           | ~64%          | \$2.50             |
| Llama 4 Maverick  | 62%                | ~78%        | N/A           | 41–54%        | Free               |
| Qwen 2.5 Coder    | N/A                | N/A         | ~31%          | N/A           | \$3                |
| Cohere Command R+ | ~88%               | N/A         | N/A           | N/A           | \$5                |

### Research methods and experimental design

This study adopts a mixed-methods research design that combines a systematic literature review, the development of a theoretical architecture, and the evaluation of an experimental prototype. Inclusion criteria required that sources address model architectures, evaluation benchmarks, or practical applications of LLM-based code generation; sources addressing hardware-description languages or non-executable code were excluded unless directly relevant to the research questions. The theoretical modelling component employed formal notation and UML-style architectural

descriptions to specify the proposed hybrid model, enabling rigorous comparison with established architectures such as AgentCoder and Self-Debug frameworks. The experimental component implemented a Python-based prototype using publicly available open-source libraries, enabling independent reproducibility of reported results.

The proposed architecture addresses three recurring failure modes identified in the literature review: partial incompleteness due to insufficient contextual grounding, factual imprecision stemming from the absence of verifiable knowledge anchors, and semantic inconsistency across independently sampled outputs. The hybrid design responds to each of these through targeted architectural components. The Retrieval-Augmented Generation component enhances generation by dynamically retrieving relevant code snippets from a curated vector database prior to each generation step. When a user submits a natural-language programming task, the system first embeds the query using QZhou-Embedding — a contextual embedding model based on the Qwen2.5-7B-Instruct architecture with a 4,096-dimensional representation space and support for sequences up to 8,000 tokens — and then queries a FAISS vector index to retrieve the top-k most semantically similar code snippets from the corpus. These retrieved snippets are prepended to the generation prompt as contextual examples, grounding the output in verified implementations and reducing the probability of hallucinated or incorrect code structures.

The Mixture of Experts component routes each generation request to a subset of specialised sub-models selected according to the detected characteristics of the prompt. Three specialist modules are defined: a Python expert module optimised for algorithmic and data-processing tasks, a JavaScript expert module for web development patterns, and a security analysis module that evaluates generated code for common vulnerability patterns. The routing mechanism uses a lightweight gating network trained on task classification signals derived from prompt features.

Self-reflection loops constitute the third architectural element. After initial code generation, the system submits the output to a structured evaluation procedure that checks for syntax errors via AST parsing, semantic alignment via GEval-style chain-of-thought reasoning, and test-case validation when reference tests are available. Outputs that fail any check are refined through additional generation steps conditioned on the identified deficiency, up to a configured maximum of 5 iterations. This mechanism directly addresses the semantic inconsistency failure mode by introducing a feedback signal that constrains subsequent outputs.

The prototype was implemented in Python 3.12.3 using the Hugging Face Transformers library (version 4.45.0) for model loading and inference, PyTorch 2.4.0 with CUDA 12.1 for GPU-accelerated computation, FAISS 1.8.0 for approximate nearest-neighbour retrieval, and NumPy 1.26.4 for vector operations. Experiments were conducted on a workstation equipped with an NVIDIA RTX 3090 GPU (24 GB VRAM), an AMD Ryzen 9 7950X processor, and 64 GB DDR5 RAM, running Ubuntu 24.04 LTS.

The code corpus used for RAG retrieval was constructed from a 10,000-snippet subset drawn from publicly available Python repositories, pre-processed to remove comments and normalise whitespace. Each snippet was embedded using QZhou-Embedding and indexed in a FAISS flat inner-product index with L2 normalisation to

enable cosine-similarity-based retrieval. Code quality was assessed using three static analysis tools: pylint for overall style and error scoring on a 0–10 scale, radon for cyclomatic complexity and Maintainability Index computation, and flake8 for PEP 8 compliance verification.

The primary evaluation dataset comprised the complete HumanEval benchmark of 164 Python programming problems. For each problem, three generations were attempted per model (pass@1, pass@5, and pass@10 metrics were approximated using standard unbiased estimators). The evaluation metric for functional correctness was pass@k, where a response is counted as correct if at least one of the k generated candidates passes all provided unit tests without modification. Average runtime per problem was recorded to assess computational efficiency. Code quality metrics were computed for all generated solutions that passed the functional correctness check, and scores were averaged across problems to produce aggregate statistics.

### Results of the Research

Characterization of the QZhou-Embedding and FAISS retrieval subsystem established the computational baseline for the hybrid architecture. Single-query embedding latency on the GPU was 0.3 milliseconds per code snippet, representing a 2.67-fold speedup over CPU execution at 0.8 milliseconds. Batch processing experiments across batch sizes from 1 to 512 revealed that GPU throughput peaked at approximately 85,000 tokens per second at a batch size of 256, after which VRAM saturation caused a modest reduction to 78,000 tokens per second at a batch size 512 (Table 6).

Table 6. Embedding Throughput by Batch Size (GPU vs. CPU)

| Batch Size | GPU Throughput (tokens/sec) | CPU Throughput (tokens/sec) | Notes                                  |
|------------|-----------------------------|-----------------------------|----------------------------------------|
| 1          | 12,000                      | 1,200                       | Baseline single-query latency          |
| 32         | 65,000                      | 6,500                       | Optimal for low-latency applications   |
| 256        | 85,000                      | 8,200                       | Peak GPU throughput                    |
| 512        | 78,000                      | 7,900                       | Minor degradation due to VRAM pressure |

Retrieval accuracy on the HumanEval subset, measured using Precision@5 and Mean Reciprocal Rank, achieved values of 0.92 and 0.88 respectively. These results confirm that QZhou-Embedding produces sufficiently discriminative code representations to support effective semantic retrieval. Retrieval latency remained below 1 millisecond for a corpus of 10,000 snippets on GPU, scaling logarithmically to 3 milliseconds for a corpus of one million snippets, consistent with the theoretical complexity of FAISS approximate nearest-neighbour search. The preprocessing step — comment removal and whitespace normalisation — produced a statistically significant 4.2 percentage point improvement in Precision@5 (paired t-test:  $t = 9.87$ ,  $p < 0.001$ ), confirming its contribution to retrieval quality. In the end-to-end pipeline, the retrieval stage accounted for less than 1% of the total query resolution time, while the generation stage consumed approximately 94% of the per-query latency. At a concurrent load of 512 simultaneous

queries, latency increased by only 18% with no measurable degradation in retrieval accuracy, indicating that the system scales acceptably under production-level concurrency. Table 7 summarises the performance of the proposed hybrid model against the GPT-4o and CodeLlama baselines on the HumanEval benchmark. The hybrid RAG-MoE model achieved a Pass@1 score of 85%, a Pass@5 score of 92%, and a Pass@10 score of 95%, with an average per-problem runtime of 320 milliseconds (Table 7).

Table 7. Performance Metrics on HumanEval Benchmark

| Model                     | Pass@1 (%) | Pass@5 (%) | Pass@10 (%) | Avg. Runtime (ms) |
|---------------------------|------------|------------|-------------|-------------------|
| Hybrid RAG-MoE (Proposed) | 85         | 92         | 95          | 320               |
| GPT-4o                    | 74         | 88         | 92          | 280               |
| CodeLlama                 | 70         | 80         | 85          | 350               |

The absolute improvement in Pass@1 over CodeLlama amounts to 15 percentage points, corresponding to approximately 24 additional correctly solved problems out of the 164 in the dataset. The improvement over GPT-4o is 11 percentage points, or approximately 18 additional problems. The self-reflection mechanism accounted for a substantial proportion of these gains: a case-study analysis of the palindrome-checking problem illustrated how an initial incorrect implementation that failed to filter non-alphanumeric characters was identified and corrected through a single reflection iteration, thereby converting a failing response into a passing one. The RAG component contributed most visibly to problems requiring recursive or algorithmically standard implementations, where retrieved examples provided accurate structural templates that the generation model could adapt. In contrast, the MoE routing mechanism showed the largest benefit for problems involving security-sensitive operations — such as input validation and boundary checking — where the security expert sub-model applied domain-specific constraints not present in the general-purpose generation path. Runtime overhead relative to CodeLlama was modest at +30 milliseconds (i.e., the hybrid model was marginally faster than CodeLlama on average), despite the additional retrieval and reflection steps. This efficiency arises because the MoE gating reduces the effective computational load per query by activating only relevant expert modules, more than compensating for the retrieval and reflection overheads. Static analysis of generated solutions across the 164 HumanEval problems produced the quality metrics reported in Table 8.

Table 8. Code Quality Metrics for Generated Solutions (HumanEval)

| Model                     | Pylint Score (/10) | Cyclomatic Complexity | Maintainability Index | Avg. LOC |
|---------------------------|--------------------|-----------------------|-----------------------|----------|
| Hybrid RAG-MoE (Proposed) | 8.5                | 2.3                   | 85                    | 12       |
| GPT-4o                    | 7.8                | 3.1                   | 78                    | 18       |
| CodeLlama                 | 7.2                | 3.5                   | 72                    | 15       |

The hybrid model's outputs received the highest pylint score of 8.5 out of 10, compared to 7.8 for GPT-4o and 7.2 for CodeLlama. Cyclomatic complexity was lowest for the hybrid model at 2.3, indicating that its solutions were structurally simpler and therefore more easily comprehensible and testable than those of the baselines. The Maintainability Index, computed as a composite of Halstead volume, cyclomatic complexity, and lines of code, similarly favoured the hybrid model at 85, relative to 78 for GPT-4o and 72 for CodeLlama. GPT-4o's relatively verbose outputs — averaging 18 lines of code against the hybrid model's 12 — account for its lower maintainability score, as the metric penalises unnecessary code length. CodeLlama's lowest quality scores reflect its tendency toward inconsistent indentation and variable naming, which the PEP 8 checker penalises independently of functional correctness. These results suggest that the RAG component's retrieval of high-quality reference implementations imposes a beneficial stylistic constraint on the generation process, in addition to its contribution to functional correctness. Despite the demonstrated improvements, the hybrid model exhibits three recurring patterns of limitations that merit targeted investigation. First, retrieval noise — cases in which retrieved snippets are syntactically similar but semantically divergent from the target task — accounted for approximately 5% of generation errors. In these cases, the retrieved context misleads the model rather than informing it, producing outputs that replicate the structure of the retrieved code without satisfying the actual task requirements. Fine-grained relevance filtering and re-ranking of retrieved candidates before prompt augmentation represent a promising direction for mitigation. Second, the self-reflection mechanism incurs approximately 14 milliseconds of runtime overhead per reflection iteration. For time-sensitive applications, this overhead may be prohibitive if multiple reflection cycles are required. Adaptive termination criteria — which halt reflection once a quality threshold is reached rather than always completing a fixed number of iterations — could reduce average overhead without sacrificing quality on the subset of problems that genuinely benefit from multiple refinement steps. Third, the current implementation supports Python and JavaScript, leaving low-resource languages such as Rust, COBOL, and niche domain-specific languages unaddressed. Extending the corpus and fine-tuning the embedding model on multilingual code corpora would be required to broaden the system's applicability to these languages, which are of significant practical importance in embedded systems, financial services, and scientific computing. At a broader level, the analysis of existing systems identifies three systemic gaps. Security evaluation, while improving with the emergence of SEC-Bench, remains immature relative to the functional correctness evaluation ecosystem. The ethical dimensions of AI code generation — including bias in training data that propagates into generated code and potential displacement of entry-level development roles — are underexplored in the empirical literature. Finally, real-world SWE-Bench resolution rates of 30–70% across all evaluated systems indicate that even the most capable current models resolve fewer than three in four real-world software engineering tasks, leaving substantial headroom for improvement.

Based on the experimental findings and the comparative analysis of existing systems, the following recommendations are offered for practitioners seeking to integrate LLM-based code generation into professional workflows. Commercial frontier models —

Claude 3.7 Sonnet and GPT-4o — provide the highest capability ceiling for complex, multi-file software engineering tasks and are most appropriate for enterprise contexts where accuracy and reliability justify their cost. Open-source alternatives — DeepSeek V3 and Qwen 2.5 Coder — offer competitive performance at substantially lower cost for function-level tasks, making them suitable for high-volume, cost-sensitive applications. The hybrid RAG-MoE approach demonstrated in this study is particularly well-suited to organisations that maintain proprietary code repositories from which high-quality retrieval examples can be drawn, and that require code outputs conforming to internal style and security standards enforceable through the MoE specialist modules. Security compliance requirements under frameworks such as GDPR and SOC 2 should be addressed through dedicated post-generation scanning rather than relying on model-level security guarantees, as current SEC-Bench results confirm persistent vulnerability rates across all evaluated models. For IDE integration, real-time latency requirements of below 100 milliseconds for code completion suggestions make smaller, locally deployed models preferable over API-dependent frontier systems, with the hybrid architecture's modular design supporting selective component activation based on task complexity.

### Conclusions

The literature review traced the evolution of code generation from rule-based synthesis through machine-learning-based approaches to the current LLM paradigm, characterised the three principal architectural families — autoregressive decoders, masked language models, and encoder-decoder models — and synthesised performance data across leading benchmarks for seven state-of-the-art systems. The comparative analysis revealed that while commercial frontier models achieve the strongest absolute performance, open-source alternatives have substantially narrowed the gap, and that no single model dominates across all benchmark dimensions simultaneously. The proposed hybrid model achieved a Pass@1 of 85% on HumanEval, exceeding CodeLlama by 15 percentage points and GPT-4o by 11 percentage points, while producing code with lower cyclomatic complexity, higher pylint scores, and higher Maintainability Index values than both baselines. The retrieval subsystem demonstrated sub-millisecond latency, high retrieval precision, and logarithmic scaling behaviour, confirming its viability as a production component.

### References

1. Jiang, J., Wang, F., Shen, J., Kim, S., & Kim, S. (2024). A survey on large language models for code generation. arXiv preprint arXiv:2406.00515. <https://arxiv.org/abs/2406.00515>
2. Lyu, M., Ray, B., Roychoudhury, A., & Tan, S. (2025). Automatic programming: Large language models and beyond. ACM. <https://dl.acm.org/doi/abs/10.1145/3708519>
3. Hou, X., Zhao, Y., Liu, Y., Yang, Z., Wang, K., & Li, L. (2024). Large language models for software engineering: A systematic literature review. ACM. <https://dl.acm.org/doi/abs/10.1145/3695988>

4. Wang, Z. (2025). Advancements and challenges of large language model-based code generation and completion. SCITEPRESS. <https://www.scitepress.org/Papers/2024/132718/132718.pdf>
5. Bistarelli, S., Fiore, M., Mercanti, I., & Mongiello, M. (2025). Usage of large language model for code generation tasks: A review. SN Computer Science. <https://link.springer.com/article/10.1007/s42979-025-04241-5>
6. Huynh, N., & Lin, B. (2025). Large language models for code generation: A comprehensive survey of challenges, techniques, evaluation, and applications. arXiv preprint arXiv:2503.01245. <https://arxiv.org/abs/2503.01245>
7. Husein, R., Aburajouh, H., & Catal, C. (2025). Large language models for code completion: A systematic literature review. Computer Standards & Interfaces. <https://www.sciencedirect.com/science/article/pii/S0920548924000862>
8. Siddiq, M., Santos, J. D. S., & Tanvir, R. (2024). Using large language models to generate JUnit tests: An empirical study. ACM. <https://dl.acm.org/doi/abs/10.1145/3661167.3661216>
9. Rasheed, Z., Waseem, M., Kemell, K., & Ahmad, A. (2025). Large language models for code generation: The practitioners perspective. arXiv preprint arXiv:2501.16998. <https://arxiv.org/abs/2501.16998>
10. Chen, L., Guo, Q., Jia, H., Zeng, Z., Wang, X., & Xu, Y. (2024). A survey on evaluating large language models in code generation tasks. arXiv preprint arXiv:2408.16498. <https://arxiv.org/abs/2408.16498>

## GAUSSIAN METHOD FOR SOLVING SLAES: PERFORMANCE COMPARISON OF MPI AND OPENMP IMPLEMENTATIONS

**Savin Yuriy**

senior lecturer

Department of Information systems

Simon Kuznets Kharkov National University of Economics, Ukraine

**Abstract.** Solving Systems of Linear Algebraic Equations (SLAEs) is a cornerstone of numerical analysis and scientific modeling. Among the various solutions available, the Gaussian method is widely used for its reliability and algorithmic clarity. As the size of the system increases, sequential execution becomes impossible, requiring the use of High-Performance Computing (HPC) technologies such as Message Passing Interface (MPI) and OpenMP. This paper presents a comparative study evaluating the performance and scalability characteristics of MPI and OpenMP implementations for the Gaussian method.

**Keywords:** Gaussian method, SLAEs, MPI, OpenMP, parallel computing, matrix computation, high-performance computing.