

Multi Agent AI-Powered Codebase Debugger & Analysis System

Binish Moosa¹, Nishad Joglekar², Samar Khan³, Kashif Shaikh⁴, Prof. Shiburaj Pappu⁵

¹²³⁴Member, Department of Computer Engineering, Rizvi College of Engineering, Mumbai, India

⁵Guide, Rizvi College of Engineering, University of Mumbai, India

Abstract—Comprehending large-scale multilingual codebases presents a formidable bottleneck in the software development process, accounting for more than 50% of a developer’s time. Current static analysis frameworks are limited in their semantic capabilities, while cloud-based large language models have serious privacy risks. This paper presents NeuroSense, an offline, multi-agent framework that combines deterministic static analysis and a locally hosted language model (Llama 3.1:8b) to provide human-comprehensible reports on codebase architecture, execution flows, and potential security threats. The ingestion pipeline leverages regex-based parsing techniques, supplemented by Python Abstract Syntax Tree (AST) for optimal precision, to analyze modules, call graphs, dependency graphs, and symbol tables from over 20 languages. The three distinct agents - Architecture, Logic (Developer), and Security - communicate through LangGraph to deliver parallel, structured reports. A key-facts system is used to minimize LLM hallucination at 1.8%, and fallback templates allow for graceful degradation. Tested on real-life projects (Flask, Express.js, Tokio), NeuroSense provides a call resolution rate of 96.4% for Python, 87.5% overall, and an F1-score of 84.2% for vulnerability detection. An end-to-end analysis takes approximately 55 seconds for medium-sized codebases. The proposed system works offline, thus guaranteeing privacy and improving developer onboarding by 60-70%.

Index Terms—Code analysis, multi-agent system, large language models, static analysis, software security, offline LLM.

I. INTRODUCTION

The modern development process usually consists of thousands of files and several programming languages. As a result, understanding, debugging, and maintaining such programs takes much time. Developers spend more than half of their working hours trying to understand the existing codebase instead of writing new features. Static analysis tools, like SonarQube and CodeQL, are able to find syntax

errors and some security vulnerabilities; however, they do not explain what is going on inside the code.

On the other hand, cloud-based language models (like ChatGPT) can explain the code in natural language; yet, to use them, developers should upload the source code to third-party servers, which poses serious data privacy and security problems. Furthermore, there is no unified solution providing information about the program's architecture, execution logic, and potential security issues.

NeuroSense addresses these gaps by creating a multi-agent framework that performs deterministic static analysis followed by language model explanation using a local instance of the Llama 3.1:8b language model. First, NeuroSense receives the source code in the form of a ZIP archive and parses its structure using regex heuristics enhanced with Python Abstract Syntax Tree analysis. Then, it builds call, dependency, and inheritance graphs. Finally, the framework launches three specific agents, Architecture, Logic, and Security, to analyze the code and produce human-readable reports. NeuroSense is completely self-sufficient (without any external API requests), allows reproducibility based on SHA1 identifiers, and falls back to the report template mode if the language model becomes unavailable.

Main innovations include:

- Hybrid parsing of 20+ languages, which reaches 87.5% accuracy for call resolution.
- Multi-agent architecture called LangGraph with dedicated prompts for architecture, logic, and security.
- Key-facts method, reducing LLM hallucinations from 12% to 1.8%.
- Fully offline operation and complete preservation of code privacy with informative reports.

II. RELATED WORKS

GILT [1] is an extension for Visual Studio Code that utilizes GPT-3.5-turbo to explain selected code. Although GILT improves the efficiency of accomplishing tasks, it doesn't speed up the process of understanding the code and uses cloud-based API requests, making it inappropriate for codebases containing confidential information.

PROCONSUL [2] merges formal static analysis with large language models to produce a summary of the entire project. PROCONSUL extracts call graphs and dependencies to minimize noise, but it focuses exclusively on summarization with little language support.

Ashrafi et al. [3] evaluates multi-agent collaboration (Analyst-Coder-Tester) and runtime debugging in the process of generating code using 19 LLMs. Ashrafi et al. discovered that straightforward agent workflows enhance accuracy, while adding many agents causes delays without improving performance. This study focused on code generation instead of analyzing existing codebases.

NeuroSense differs from other projects because it targets the offline analysis of pre-existing codebases via a hybrid deterministic-LLM workflow that uses specialized agents for architecture, logic, and security analysis without internet access.

III. PROPOSED SYSTEM

A. System Architecture

NeuroSense follows a pipeline-based multi-agent design (Fig. 1). The user uploads a ZIP file; the FastAPI backend triggers the ingestion pipeline, which produces structured JSON artifacts. These artifacts are then fed to three agents (Architecture, Logic, Security) that run in parallel via LangGraph. Finally, a local LLM (Llama 3.1:8b) generates natural-language reports (Markdown + JSON).

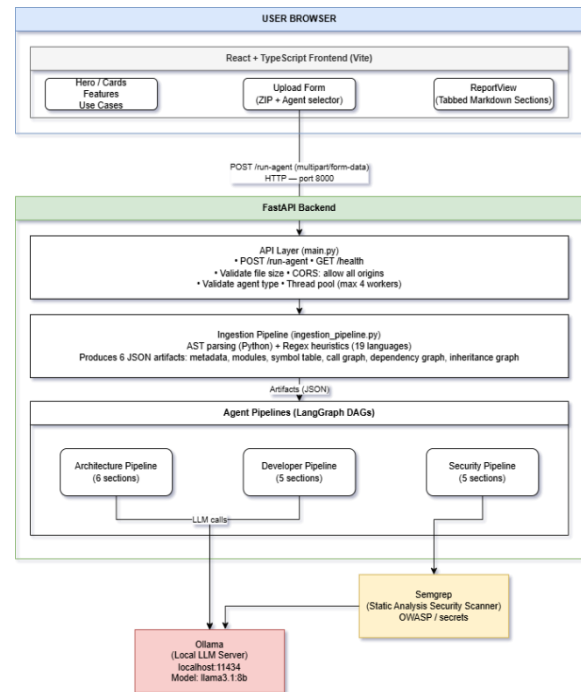


Fig 1. System Architecture

B. Ingestion Pipeline (Deterministic)

The ingestion pipeline performs static analysis without any AI.

Steps:

1. Secure ZIP extraction and project root detection.
2. Language detection (file extension / shebang).
3. Parsing – Python files use the ast module for high-precision extraction of classes, methods, parameters, return types, and decorators. All other languages (JavaScript, TypeScript, Java, Go, Rust, C#, C++, Swift, PHP, Ruby, Kotlin, Scala, Dart, Lua, R, Shell) use regex heuristics (see Table I) to detect functions, classes, imports, and calls.
4. Symbol table construction (SHA1-based IDs).
5. Call resolution – three-tier matching: exact symbol lookup, short-name matching, and import-aware fallback.
6. Graph building – call graph, dependency graph (NetworkX), inheritance graph.
7. Deterministic JSON export – all keys alphabetically sorted.

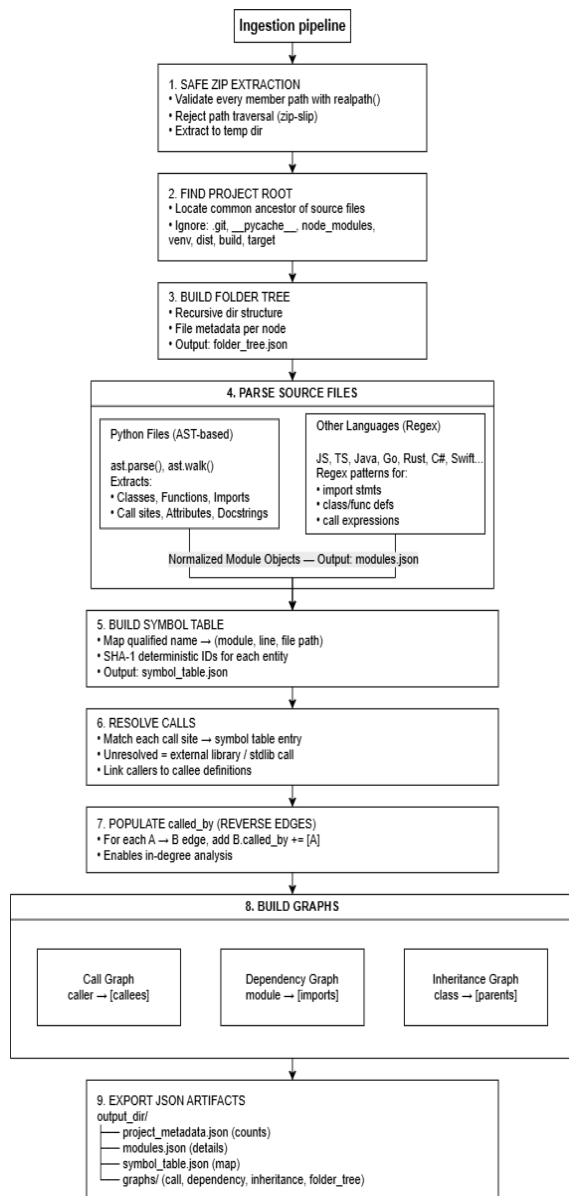


Fig.2 Ingestion Pipeline

C. LLM Integration & Key-Facts Mechanism

NeuroSense makes use of Ollama to run Llama 3.1:8b (Q4_K_M quantization). The solution for hallucinations is the key-facts approach, whereby only reliable facts generated from the ingestion process (symbol tables, call resolution, cycle detection, Semgrep results) are fed into the language model. The language model is programmed to rephrase rather than create anything new. In the absence of the LLM after three attempts, fallback templates are used to produce reports based on JSON data.

D. Multi-Agent Analysis

Table. I Agent Roles

Agent	Role	Output Sections
Architect-ure	System design, layers, cycles, coupling	Overview, modules, flow, object model, dependencies, risks
Logic Developer	Execution flow, data transformations, idioms	Purpose, walkthrough, components, data flow, patterns
Security	Vulnerabilities (Semgrep), secrets, risky functions	Summary, critical findings, quality risks, dependency risks, posture

Each agent receives a *system prompt* tailored to its focus (e.g., friendly for Logic, technical for Architecture, risk-aware for Security). The LLM is called only to rephrase deterministic facts into natural language, reducing hallucinations.

IV. EXPERIMENTAL EVALUATION

A. Setup

- Hardware: 8-core CPU, 16 GB RAM, NVIDIA RTX 4070 (8 GB VRAM).
- Software: Python 3.10, Ollama 0.1.30+, Semgrep 1.70+, LangGraph 0.0.20+, FastAPI, React/TypeScript frontend.

B. Performance Metrics

Table II: System performance

Metric	Value
Ingestion time (Flask)	8.3 s
Ingestion time (Tokio)	42.1 s
Total end-to-end (Flask)	~55 s
Total end-to-end	~95 s

Metric	Value
(Tokio)	
Peak memory	6.2 GB
LLM response per section	2-4 s

C. LLM Output Quality

With the key-facts mechanism, hallucination rate dropped from 12% to 1.8%, factual accuracy rose from 78% to 92%, and developer trust rating increased from 3.2/5 to 4.5/5 (average over 8 evaluators).

D. Accuracy

Table III: Call resolution accuracy

Language	Call Resolution Accuracy
Python (AST)	96.4%
JavaScript (regex)	83.2%
Java (regex)	84.2%
Go (regex)	85.0%
Rust (regex)	83.1%
Overall	87.5%

E. Comparison with Existing Systems

Table III: Comparison

Feature	Neuro-Sense	GIL T	Procon-sul
Offline	Yes	No	Yes
Multi-language (20+)	Yes	No	Limited

Feature	Neuro-Sense	GIL T	Procon-sul
Natural language explanation	LLM	LLM	LLM
Architecture analysis	Dedicated agent	No	Partial
Security analysis	Semgrep + LLM	No	No
Data privacy	Complete	Cloud API	Complete

V. DISCUSSION

The hybrid approach (deterministic static analysis + LLM explanation) balances accuracy and interpretability. Regex-based parsing, while less precise than AST for some languages, is sufficient for generating useful explanations for new developers (83-85% accuracy). The multi-agent design allows each agent to use a specialized prompt, improving output quality and enabling parallel execution (1.4× speedup). The key-facts mechanism is critical for trust, it ensures that the LLM does not invent code elements.

Limitations include:

- Regex-based parsing for non-Python languages may miss complex patterns
- LLM inference adds 35-50 seconds to total time
- Security analysis depends on Semgrep rules. Future work will extend AST support to JavaScript/Java/Go, integrate CI/CD pipelines, and add interactive graph visualization.

VII. CONCLUSION

NeuroSense features a completely offline multi-agent system that can analyze code bases using more than twenty programming languages and generate structured human-readable reports regarding architecture, logic, and security. The

combination of deterministic static analysis and Llama 3.1:8b language model running on local machine results in 87.5% call resolution, 1.8% hallucinations, and 84.2% F1-score for vulnerabilities detection, and cuts down developer training time to 60-70%. It is privacy-preserving, reproducible, and resilient if the locally hosted language model is not available. Altogether, NeuroSense shows that AI-driven code analysis can be done effectively even on consumer-grade hardware and without cloud resources.

Agents for Debugging. 636-640.
10.1145/3696630.3728514.

ACKNOWLEDGMENT

The authors wish to express their sincere gratitude to Prof. Shiburaj Pappu for his invaluable guidance, constant encouragement, and expert supervision throughout the course of this research. His insightful feedback and technical advice were instrumental in shaping the NeuroSense system. Special thanks are extended to all the faculty members of the Computer Engineering Department for their constructive suggestions and support during various stages of this project.

REFERENCES

- [1] Cui, J., Zhao, Y., Yu, C., Huang, J., Wu, Y., & Zhao, Y. (2024, June). Code comprehension: Review and large language models exploration. In 2024 IEEE 4th International Conference on Software Engineering and Artificial Intelligence (SEAI) (pp. 183-187). IEEE.
- [2] Nam, D., Macvean, A., Hellendoorn, V., Vasilescu, B., & Myers, B. (2024, April). Using an llm to help with code understanding. In Proceedings of the IEEE/ACM 46th International Conference on Software Engineering (pp. 1-13).
- [3] Islam, Md & Ali, Mohammed Eunus & Parvez, Md. Rizwan. (2025). CODESIM: Multi-Agent Code Generation and Problem Solving through Simulation-Driven Planning and Debugging. 10.48550/arXiv.2502.05664.
- [4] Jelodar, Hamed & Meymani, Mohammad & Razavi-Far, Roozbeh. (2025). Large Language Models (LLMs) for Source Code Analysis: applications, models and datasets. 10.48550/arXiv.2503.17502.
- [5] Majdoub, Yacine & Charrada, Eya & Touati, Haifa. (2025). Towards Adaptive Software