

Neurosymbolic Code Auditing

Suggested Citation: Chengpeng Wang, Zhuo Zhang, Xiangzhe Xu, Mingwei Zheng, Guannan Wei and Xiangyu Zhang (2026), "Neurosymbolic Code Auditing", : Vol. xx, No. xx, pp 1–18. DOI: 10.1561/XXXXXXXXXX.

Chengpeng Wang

Purdue University
wang6590@purdue.edu

Zhuo Zhang

Columbia University
zz@cs.columbia.edu

Xiangzhe Xu

Purdue University
xu1415@purdue.edu

Mingwei Zheng

Purdue University
zheng618@purdue.edu

Guannan Wei

Tufts University
guannan.wei@tufts.edu

Xiangyu Zhang

Purdue University
xyzhang@cs.purdue.edu

This article may be used only for the purpose of research, teaching, and/or private study. Commercial use or systematic downloading (by robots or other automatic processes) is prohibited without explicit Publisher approval.

now

the essence of knowledge

Boston — Delft

Contents

1	Introduction	3
1.1	Code Auditing in the AI Era	4
1.2	LLMs for Code Auditing	6
1.3	Scope and Organization	12
2	Code Auditing and Its Challenges	16
2.1	Bug Categorization	16
2.2	Static Analysis and Its Challenges	23
2.3	LLM-driven Code Analysis and Its Challenges	31
2.4	Summary	36
3	Code Auditing From an Agentic Perspective	38
3.1	Environment of Code Auditing	39
3.2	Tools for Code Auditing	40
3.3	Memory for Code Auditing	44
3.4	Summary	46
4	LLM-driven Primitive Analysis	48
4.1	Primitive Analysis over Code Artifacts	48
4.2	Primitive Analysis over Non-Code Artifacts	60
4.3	Summary	69

5 Auditing Code as Human Auditors	71
5.1 Simulating Manual Code Auditing	72
5.2 Simulating Semi-Automatic Code Auditing	86
5.3 Industrial Advances in Neurosymbolic Code Auditing	96
5.4 Summary	100
6 Conclusion and Future Directions	104
6.1 Auditable Software and Ecosystem Boundaries	106
6.2 Cost-Aware Neurosymbolic Workflows	108
6.3 Realistic Evaluation and Auditor Robustness	112
6.4 Summary	114
References	116

Neurosymbolic Code Auditing

Chengpeng Wang¹, Zhuo Zhang², Xiangzhe Xu³, Mingwei Zheng⁴,
Guannan Wei⁵ and Xiangyu Zhang⁶

¹*Purdue University; wang6590@purdue.edu*

²*Columbia University; zz@cs.columbia.edu*

³*Purdue University; xu1415@purdue.edu*

⁴*Purdue University; zheng618@purdue.edu*

⁵*Tufts University; guannan.wei@tufts.edu*

⁶*Purdue University; xyzhang@purdue.edu*

ABSTRACT

Code auditing examines whether an implementation satisfies target properties such as reliability, security, performance, and design intent. Classical static analysis provides rigorous foundations for automated auditing, but it remains difficult to apply broadly in real-world software. The emergence of LLMs and agents creates new opportunities for leveraging semantic cues that are embedded in source code, documentation, and other software artifacts yet difficult for classical static analyzers to derive.

This survey presents a comprehensive study of *neurosymbolic code auditing*, an emerging synthesis that integrates LLM-driven reasoning with classical static analysis and other structured program reasoning techniques. By formalizing the cognitive processes underlying manual code auditing, we conceptualize an agent-centric framework that serves as a unifying foundation for neurosymbolic analysis. Specifically, we categorize LLM-enabled primitive analyses over both code and non-code artifacts, addressing the fundamental

reasoning tasks routinely performed by human auditors. Furthermore, we review a broad spectrum of neurosymbolic architectures that combine symbolic tools with LLM-based reasoning, illustrating how these systems replicate and extend the practical strategies employed by human auditors. Lastly, we synthesize the survey around three recurring bottlenecks, namely observation, verification, and memory, then outline future directions and the persistent technical challenges.

1

Introduction

Code auditing is a fundamental activity in software development in which developers, security analysts, and automated tools examine whether an implementation satisfies a set of expected properties. These properties typically concern security, reliability, performance, or compliance with project-specific specifications. Rather than merely inspecting code, auditing is a multi-step, often iterative, evidence-seeking process that involves gathering program facts from source code and surrounding artifacts and assessing whether those facts indicate violations of the target properties.

From the perspective of automated support, code auditing is commonly assisted by dynamic and static program-analysis techniques. Dynamic techniques, such as unit testing (Xie *et al.*, 2016; Lukasczyk and Fraser, 2022) and fuzzing (Gutmann, 2016; Zhu *et al.*, 2022), execute the program to expose concrete failures. They can be effective at uncovering certain classes of defects, but their applicability is constrained by the availability of an execution environment, adequate test harnesses, and sufficient input coverage. Static techniques, in contrast, analyze code without executing it and can therefore be applied earlier in the development process. In this survey, we focus on *static code auditing*,

This focus is particularly relevant in modern software development, where large codebases, complex dependencies, and AI-generated code increasingly create situations in which execution-based validation is impractical, but early and scalable auditing support is still required.

This chapter first provides preliminaries on code auditing, with particular attention to the new challenges that have emerged in the AI era. It then examines the opportunities afforded by large language models (LLMs) (Brown *et al.*, 2020; OpenAI, 2023) for code auditing, discussing their intrinsic capabilities and limitations. Finally, it introduces the scope of this survey and presents an overview of its organization.

1.1 Code Auditing in the AI Era

The rise of generative artificial intelligence (Vaswani *et al.*, 2017; Brown *et al.*, 2020; Ouyang *et al.*, 2022) has fundamentally reshaped modern software development (Chen *et al.*, 2021; Jiang *et al.*, 2026). AI coding agents powered by LLMs, such as Codex (OpenAI, 2026a) and Claude Code (Anthropic, 2026a), have been rapidly evolving and are now deeply embedded in mainstream development environments and workflows. Developers increasingly rely on these tools to generate program lines, complete functions, and even implement non-trivial program logic, substantially accelerating the pace of software production (Peng *et al.*, 2023; Ziegler *et al.*, 2022; Barke *et al.*, 2023).

However, AI coding agents also raise concerns that go beyond security, including reliability, performance, maintainability, and broader aspects of code quality. For example, recent studies have demonstrated that AI-generated code can harbor substantial security weaknesses (Xu *et al.*, 2025b; Fu *et al.*, 2025). Developer surveys further indicate that practitioners frequently expend considerable effort correcting “almost-right” outputs and remain skeptical of their correctness. As a result, code auditing has taken on heightened importance in the AI era. Because developers may possess only a partial understanding of generated implementations and the design decisions behind them, they must invest greater effort in defect detection, root-cause localization, and repair. Auditing therefore serves as a critical mechanism for establishing trust in AI programming (Roychoudhury *et al.*, 2026).

Over the past several decades, a broad spectrum of program-analysis techniques has been developed to support code auditing, including data-flow analysis (Kildall, 1973; Kam and Ullman, 1977; Reps *et al.*, 1995; Arzt *et al.*, 2014), abstract interpretation (Cousot and Cousot, 1977a; Cousot and Cousot, 1979; Blanchet *et al.*, 2003), symbolic execution (King, 1976; Cadar *et al.*, 2008), and model checking (Clarke and Emerson, 1982; Queille and Sifakis, 1982; Holzmann, 1997). These static analysis frameworks offer rigorous foundations for reasoning about program behavior and have achieved considerable success. However, fundamental gaps still remain that prevent them from serving as general-purpose auditing solutions, particularly in the context of AI-assisted software development.

First, AI-assisted coding workflows often produce intermediate code artifacts that are incomplete with respect to their project contexts. For example, developers may receive partial implementations that lack surrounding type definitions, imported dependencies, build configurations, or the caller and callee context needed to construct a complete analysis target. Many production static analyzers assume a well-formed and buildable project. Although modular and compositional analyses relax the need for whole-program availability, they still typically require explicit modules that may be absent in early AI-assisted coding scenarios. This limitation is not unique to AI-assisted workflows. Even in conventional software development, existing static analysis tools generally offer limited support when developers need to reason about in-progress code that is not yet integrated into a buildable project.

Second, many mature static analyzers are most effective when the target property can be stated as a relatively precise and recurring pattern, such as memory-safety violations and taint-style bugs. These classic bug categories remain important, but the quality concerns arising in AI-assisted programming are broader and more diverse. They include subtle logical errors, semantic inconsistencies with developer intent, violations of implicit design conventions, and performance anti-patterns. The difficulty is not that such properties are inexpressible in formal languages. Rather, their intended specifications are often unavailable, underspecified, or highly context-dependent. As a result, building sound and complete characterizations of these properties of-

ten becomes impractical, while natural-language descriptions remain accessible because developers can tolerate a certain degree of ambiguity when communicating intent.

These limitations reveal a mismatch between the assumptions of classical static analysis and the demands of modern AI-assisted software development. As AI coding agents assume an increasingly large role in software production, there is a pressing need for auditing techniques that can reason about partial code artifacts, recover or approximate missing context, and accommodate properties whose specifications are distributed across diverse software artifacts.

1.2 LLMs for Code Auditing

Beyond code generation, modern LLMs and AI coding agents are increasingly employed as assistants capable of inspecting, reasoning about, and acting upon complex software artifacts within integrated development environments (IDEs). Systems such as Codex (OpenAI, 2026a), Claude Code (Anthropic, 2026a), and other AI coding agents have demonstrated that LLMs, when equipped with appropriate tooling, can navigate multi-file codebases, follow cross-module call chains, interpret documentation, and respond to high-level auditing requests expressed in natural language. This emerging paradigm points to a fundamental shift in how code auditing may be conducted. Rather than relying exclusively on classical static analysis frameworks, auditing can increasingly draw on customizable, interactive, and goal-driven semantic reasoning over both code and its surrounding artifacts.

To illustrate this potential, we consider the motivating example in Figure 1.1, a simplified null pointer dereference bug drawn from `sofa-pbrpc`¹, an open-source RPC framework. The bug involves a null value that originates in the function `field2json`, escapes the function boundary, and is eventually dereferenced without a guard. When we provide Claude 3.5 Sonnet with the set of relevant functions and specify, in natural language, the auditing objective of detecting the null pointer dereference, the model is able to analyze the code and produce a

¹<https://github.com/baidu/sofa-pbrpc>

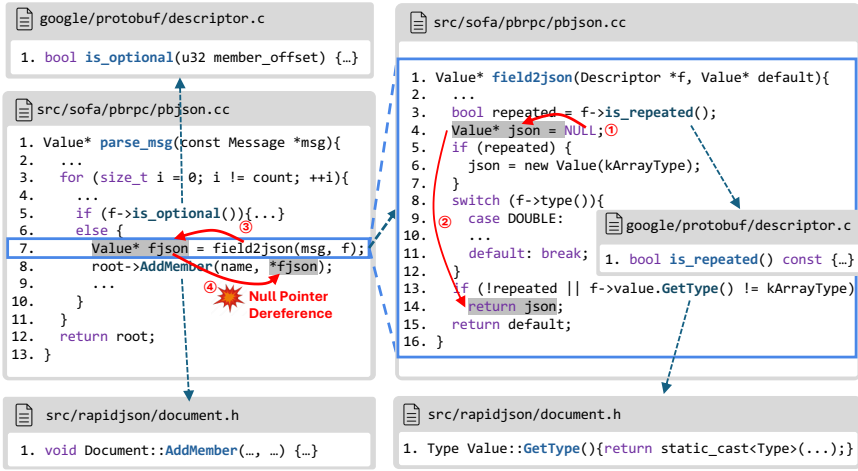


Figure 1.1: A simplified null pointer dereference bug drawn from the open-source RPC framework `sofa-pbrpc`. The function `field2json` may return null under certain branch conditions. The null value is propagated out of the function `field2json` before being dereferenced without a null check. Blue dashed arrows indicate call-graph edges, while red solid arrows trace the interprocedural data-flow path along which the null value propagates from its origin to the dereference site.

detailed explanation of how the null value is propagated through the call chain and eventually dereferenced. Notably, the model does not require constructing a complete intermediate representation of the project, resolving external dependencies, or writing formal specifications. Instead, the auditing intent can be specified directly in natural language, and the model responds by interpreting the code semantics, tracing data-flow facts, and generating a detailed bug explanation.

This interaction exemplifies several important departures from classical static analysis. First, auditing can be brought earlier into the development cycle, for example directly inside an IDE, without requiring a fully configured analysis environment or a successful build. Second, the auditing objective can be specified flexibly. Human auditors can ask about particular bug classes, project-specific invariants, or ad hoc concerns without implementing new checkers. Third, the output can take the form of an explanation rather than a bare warning, connecting root causes, propagation paths, and failure points in a manner similar to human audit reports. These properties suggest that LLMs can serve

as powerful building blocks for next-generation code-auditing systems, offering flexibility, accessibility, and semantic expressiveness that are difficult to obtain from classical static analysis alone. At the same time, LLMs remain insufficient as standalone auditors. In the following sections, we discuss both their intrinsic capabilities and the limitations that prevent them from scaling reliably to large-scale code auditing.

1.2.1 LLMs’ Intrinsic Capabilities

Prior studies (Chen *et al.*, 2024; Cheng *et al.*, 2024; Yadavally *et al.*, 2024) have shown that LLMs can perform basic static-analysis tasks on small code fragments. This capability can complement classical static analysis, which typically operates over structured intermediate representations, such as Static Single Assignment (SSA) (Cytron *et al.*, 1991), and often requires substantial engineering to model program semantics, such as points-to properties and data dependencies. In contrast, LLMs can often conduct such reasoning within a limited analysis scope by leveraging code-understanding abilities acquired during pretraining and post-training, much like skilled human auditors. In what follows, we focus on three critical capabilities, including *reducing programs*, *handling pointers*, and *exploring feasible program paths*.

Capability I: Reducing Programs. Program reduction (Ferles *et al.*, 2017; Sun *et al.*, 2018) is essential for scalability in code auditing. Given a property, a set of initial program points, and a defined scope, program reduction identifies the subset of statements relevant to the property within the scope. These statements form a self-contained, smaller program, significantly reducing the number of paths to analyze. For instance, in null pointer dereference detection, program reduction targets how null values are propagated in the program and focuses on key statements such as null pointer assignments, conditional checks guarding propagation, and cross-function propagation (e.g., passing the pointer to a callee function via a parameter, returning it to a caller function, or writing it to a global variable). In a preliminary experiment (Figure 1.2), we fed the function `field2json` shown in Figure 1.1 to Claude 3.5 Sonnet and asked it to reduce the program with respect to

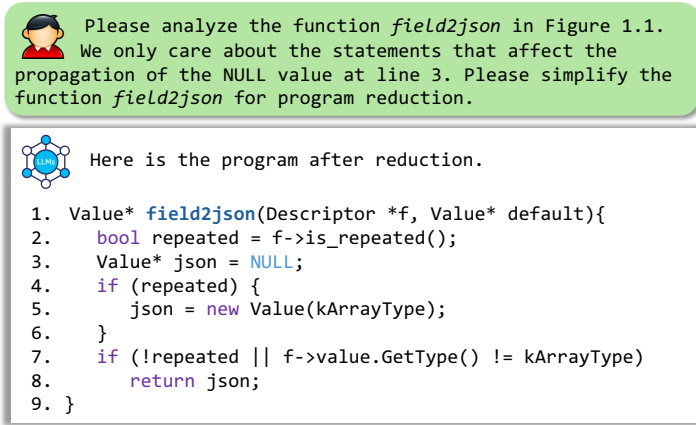


Figure 1.2: An example of program reduction by LLMs. Given the function `field2json` from Figure 1.1 and the targeted null value at line 3, the model (Claude 3.5 Sonnet) retains only the statements relevant to how a null value is produced and returned, while removing irrelevant logic such as the `switch` block (lines 8–12 in the function `field2json`).

null-value propagation. The reduced program retains only the statements critical to how a null value is produced and returned, while removing irrelevant logic such as the `switch` block. Notably, human auditors often perform such reduction implicitly, facilitating complex code reasoning without the excessive path exploration that hampers classic static tools.

Capability II: Handling Pointers. Program pointers may point to different memory locations depending on their runtime values. Consequently, reads and writes via pointer dereferences create data-flow facts that dynamically vary based on pointer values. Determining the set of memory locations a pointer variable may point to is known as *points-to analysis* (Smaragdakis and Balatsouras, 2015), an important prerequisite of many downstream analyses, such as data-dependence analysis. Classical points-to analysis techniques are based on memory abstractions, treating concrete memory locations as abstract objects. For instance, an analysis may assume that a pointer referencing a heap object could potentially point to multiple heap objects, as precise heap addresses cannot typically be determined statically. This over-approximation leads to overly inflated dependence relations, complicating downstream analysis

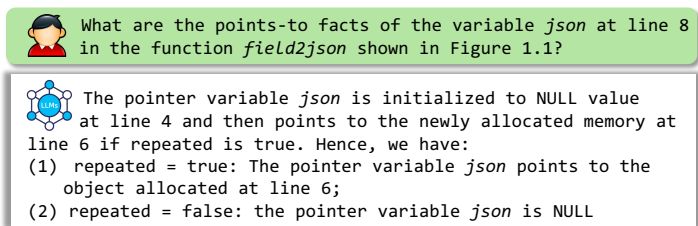


Figure 1.3: An example of pointer handling by LLMs. Given `field2json` shown in Figure 1.1, the model (Claude 3.5 Sonnet) identifies the two possible points-to targets of the return value, namely a heap-allocated JSON object and null, together with the path conditions under which each arises.

tasks. In contrast, human auditors can often intuitively and accurately determine the points-to facts of a specific pointer variable through their understanding of program semantics. Advanced LLMs exhibit similar capabilities, particularly within the scope of individual functions. As shown in Figure 1.3, we prompt Claude 3.5 Sonnet with the `field2json` function from Figure 1.1 and query the points-to targets of its return value. The model accurately identifies the two possible targets, including a heap object allocated at line 6 in the function `field2json` and null, and, crucially, reports the path conditions under which each arises. In contrast, representative static analyzers such as SVF (Sui and Xue, 2016) compute points-to facts without path conditions by default, conflating the two cases and losing precision.

Capability III: Exploring Feasible Program Paths. A bug report is often considered valid only if it provides a feasible program execution path from the root cause to the symptom as evidence. To determine feasibility, classical tools rely on modeling conditional checks along a path as symbolic constraints (e.g., in the form of first-order logic formulas) and querying a theorem prover (e.g., the SMT solver Z3 (Moura and Bjørner, 2008)) to check if an input can satisfy the constraints. This process is computationally expensive and prone to failure, as converting program paths to logic formulas requires exploring an exponentially large number of program paths and modeling a wide range of program behaviors that lack direct logical representations, such as loops, recursion, aliasing, pointer arithmetic, and unbounded string operations.

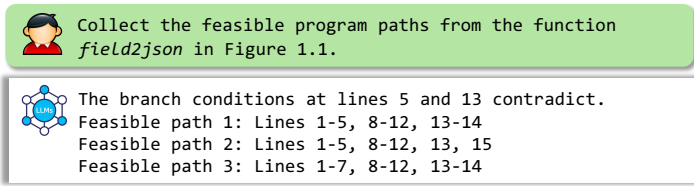


Figure 1.4: An example of feasible program path exploration by LLMs. Given `field2json` shown in Figure 1.1, the model (Claude 3.5 Sonnet) skips irrelevant branches (e.g., the `switch` block), detects that the conditions at lines 5 and 13 are contradictory, and enumerates the three feasible paths through the function.

In contrast, humans, as well as LLMs, rely on program reduction and intuitive logical reasoning to assess feasibility, an approach that is highly effective within a limited scope. For example, in Figure 1.4, we prompt Claude 3.5 Sonnet to enumerate the feasible paths through `field2json` from Figure 1.1. The model skips the irrelevant `switch` block, discovers that the branch conditions at lines 5 and 13 are contradictory, thereby ruling out paths that traverse both lines 6 and 14, and reports exactly three feasible paths. In contrast, a classical symbolic execution engine would need to encode every branch as a solver query, incurring substantially higher cost on real-world code with deeply nested control flow.

1.2.2 LLMs’ Inadequacy

LLMs alone are not silver bullets for code auditing. Due to the inherent complexity of software engineering (Brooks, 1987) and the undecidability of static analysis problems (Rice, 1953), we cannot expect LLMs to discover all bugs in real-world codebases without false positives. Recent frontier models have substantially expanded the usable context window (Liu *et al.*, 2024a; Chen *et al.*, 2023; Peng *et al.*, 2024), making it possible to process larger code fragments. However, a large nominal context window does not by itself guarantee reliable code auditing. Reasoning over long contexts is still constrained by finite test-time compute and by the model’s imperfect ability to retrieve, connect, and reason over relevant information scattered across the input. Notably, recent long-context evaluations show that model performance can vary

substantially across downstream tasks and may degrade as the context becomes longer (Yen *et al.*, 2024). Therefore, long-context modeling should be viewed as a useful substrate for code auditing rather than a complete solution.

To validate our speculation, we conducted a controlled experiment in which we prompted Claude 3.5 Sonnet with all five functions shown in Figure 1.1 and asked it to identify the null pointer dereference bug in a single end-to-end pass, following a methodology similar to a recent study (Fang *et al.*, 2024a). We call this a controlled experiment because, in practice, an auditor cannot assume knowledge of the complete set of functions relevant to a bug. The model exhibited substantial hallucinations, reporting that almost all dereferenced pointers have null values. Even when we improved the prompts by offering several few-shot examples and explanations of how a null pointer dereference bug occurs, the model still hallucinated, producing false positives and incorrect explanations. Therefore, an end-to-end prompting-based solution cannot offer high-quality auditing results.

This gap reflects a well-known structural challenge in manual code auditing: given scale, complexity, and limited time, even experienced human auditors cannot reliably identify all potential bugs in a large and complex software system through direct codebase inspection alone. This observation motivates the exploration of a more integrated approach to code auditing. In this survey, we refer to the systematic orchestration of LLMs and static analysis techniques within auditing workflows as *neurosymbolic code auditing*, an emerging synthesis that constitutes the central theme of this work.

1.3 Scope and Organization

This survey adopts a neurosymbolic perspective on code auditing. Rather than treating auditing as a single-shot prediction problem or a purely symbolic reasoning task, we view it as an evidence-seeking workflow. In this workflow, an auditor operates over a rich environment of code and non-code artifacts, interacts with external static-analysis tools, and incrementally accumulates the program facts needed for concrete auditing tasks. This perspective provides a unifying lens for recent

code-auditing techniques, clarifies the main bottlenecks they face, and delineates the distinct roles of different tool families.

Although the neurosymbolic auditing framework introduced here is, in principle, applicable to a broad range of coding tasks, including code understanding, refactoring, and transformation, this survey focuses primarily on *bug detection*, especially defects relevant to security and reliability. Specifically, bug detection is the application domain in which the convergence of LLMs and static-analysis techniques has been most actively explored, and it provides the concrete technical challenges around which the surveyed literature is organized. Although a small number of surveyed techniques incorporate dynamic components, such as concolic execution or hybrid fuzzing (Meng *et al.*, 2024a; Luo *et al.*, 2026), their primary reasoning remains static in nature. We include these techniques because they instantiate the same neurosymbolic mechanisms that unify this survey.

Under this lens, this survey covers work published between 2023 and 2026 across the programming languages, software engineering, computer security, machine learning, and natural language processing communities. For programming languages, software engineering, and computer security venues, we collected bibliographic metadata from conference proceedings and used Claude Code with Claude Sonnet 4.6 to screen paper abstracts for relevance to LLMs and code auditing. Because the machine learning and natural language processing literatures are substantially larger, we relied on citation-guided exploration in these communities and selected representative and influential works. We also include selected high-quality arXiv papers when they illuminate important emerging directions.

To ground the survey in the foundations of static analysis, we begin with representative works across major theoretical frameworks and application domains. Their limitations motivate our investigation of neurosymbolic approaches that combine classical static-analysis techniques with LLM capabilities. Overall, our goal is to provide broad coverage of both recent neurosymbolic code-auditing research and the classical static-analysis literature, while acknowledging that several relevant work may nevertheless have been omitted. The raw data collected for the survey, including venue `.bib` files, HTML pages containing accepted-paper

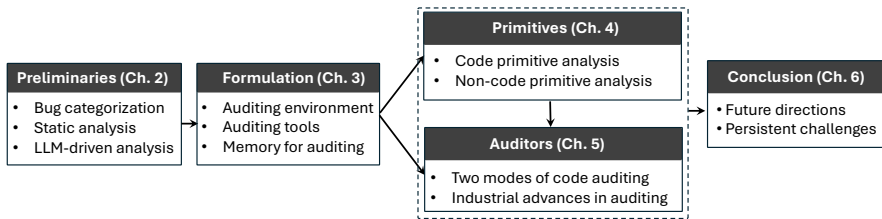


Figure 1.5: The organization of the survey

lists, and paper-classification scripts, are available online².

The remainder of this survey is organized as a progressive argument, which is demonstrated by Figure 1.5. Chapter 2 establishes the problem setting by revisiting the foundations of code auditing and common bug categories, examining why formally rigorous static-analysis techniques remain brittle in practice, and explaining why LLM-only approaches are insufficiently reliable for end-to-end auditing. Chapter 3 then formalizes code auditing from an *agentic* perspective, where an auditor analyzes code and non-code artifacts in the environment, augments its observations through tools such as static analyzers, and relies on memory to track evolving hypotheses, partial facts, and verification traces across codebase navigation and task decomposition. Building on this abstraction, Chapter 4 examines how LLMs can perform *primitive analyses* that extract primitive program properties and interpret non-code artifacts, positioning them as bounded-scope components rather than wholesale replacements for classical analyses. Chapter 5 shifts from primitive components to complete systems by surveying neurosymbolic auditing workflows that support or automate key aspects of human auditing. These include systems that simulate *manual auditing* by using LLMs to interpret program semantics, as well as *semi-automatic auditing* in which LLMs help customize and configure static analyzers and interpret their results. Finally, Chapter 6 synthesizes the survey around three coupled bottlenecks, including *observation*, *verification*, and *memory*, and reframes neurosymbolic auditing as an engineering discipline aimed at repeatedly producing *stable*, *actionable* security conclusions under real-world constraints. It closes by outlining future research directions,

²<https://github.com/PurCL/ASE>

including auditability-by-design interfaces, cost-aware orchestration, evidence-prioritized remediation loops, ecosystem-level auditing across configurations and dependencies, the translation of neural reasoning into reusable artifacts with explicit lifecycles and uncertainty, and the security of the auditor itself in adversarial environments.

2

Code Auditing and Its Challenges

In this survey, we consider bug detection the primary application scenario for code auditing, where the audited properties facilitate security analysis and inform subsequent actionable protection measures. To establish the necessary background, we first classify bugs according to the taxonomy defined by the Common Weakness Enumeration (CWE, 2025), whose categories have been extensively studied by the research community over the past several decades. We then revisit existing code auditing techniques, encompassing both classical static analysis and LLM-driven analysis, and examine their respective strengths and limitations. Finally, we introduce the key insights underlying neurosymbolic code auditing, which motivate the remaining chapters of this survey.

2.1 Bug Categorization

Bug categorization provides a fundamental lens for understanding both the challenges of code auditing and the strategies needed to address them. By categorizing bugs according to the relevant program expressions they involve and the associated semantic properties, we identify four major bug classes, namely numeric bugs, memory bugs, taint-style bugs, and functional bugs. This taxonomy highlights the diversity of reasoning

tasks involved in effective auditing and provides a basis for assessing how well existing techniques support different bug classes.

2.1.1 Numeric Bugs

Numeric bugs arise when arithmetic computations operate outside valid numeric domains or violate developer-intended numeric contracts, leading to exceptions, incorrect results, or security-relevant program states. Representative examples include divide-by-zero errors (Guo *et al.*, 2022), out-of-range conversions, integer overflow and underflow (Pomonis *et al.*, 2014; Dietz *et al.*, 2012), and floating-point errors. Floating-point bugs are especially subtle because floating-point numbers only approximate real values with finite precision. As a result, rounding errors and error amplification can cause computations to deviate from the intended mathematical behavior, as observed in machine learning systems (Zhang *et al.*, 2020). The core reasoning task in numeric bug detection is therefore to determine whether program values remain within valid and intended numeric ranges across arithmetic operations, type conversions, and control-flow paths.

Notably, numeric anomalies such as overflow and underflow are not intrinsically erroneous (Dietz *et al.*, 2012). Their status depends on both language-defined semantics and the developer’s intent. For instance, signed integer overflow may constitute undefined behavior in several programming languages, whereas unsigned overflow or modular arithmetic may be well defined in others. Even when wrap-around behavior is well defined, it may be either intentional, as in hashing and checksum computations, or unintended, as in allocation-size calculations and bounds checks. Consequently, the mere presence of a numeric anomaly is insufficient to establish a bug. The anomaly must be interpreted in relation to the applicable execution semantics and the numeric contract expected by the surrounding program.

For these reasons, detecting numeric bugs requires reasoning that is both *semantics-aware* and *intent-aware*. On the semantics side, an auditor must track value ranges, path conditions, and floating-point error propagation with sufficient precision to determine whether an operation may exceed representable bounds, trigger exceptional behavior, or

accumulate round-off error relative to the intended real-number computation. Classical analyses address these checks using numerical abstract domains, such as affine equalities, intervals, polyhedra, octagons, and related relational domains (Miné, 2002; Miné, 2017; Jeannet and Miné, 2009), or using symbolic execution with constraint solving (Godefroid *et al.*, 2005; Sen *et al.*, 2005; Cadar *et al.*, 2008). For floating-point programs, the analysis must additionally account for the gap between finite-precision execution and the idealized real-number computation, including rounding, cancellation, and error amplification (Monniaux, 2008; Goubault and Putot, 2011; Darulova and Kuncak, 2014). On the intent side, the auditor often needs contextual evidence, such as type annotations, surrounding bit-level idioms, and other specifications, to determine whether wrap-around, truncation, or precision loss is intentional and acceptable. However, most existing code auditing techniques focus primarily on language-defined numeric semantics, flagging potential overflows or exceptional cases on the basis of abstract execution alone, while largely neglecting intent-aware distinctions. As a result, these techniques tend to over-report intentional computations.

2.1.2 Memory Bugs

Memory bugs arise when programs misuse memory objects or violate the spatial, temporal, or lifetime assumptions required for safe memory management. They are a major source of reliability failures and security vulnerabilities in low-level programs, especially C/C++ systems software. Representative examples include null-pointer dereferences (Xie *et al.*, 2003), spatial errors such as out-of-bounds reads and writes (Austin *et al.*, 1994; Nagarakatte *et al.*, 2009), temporal errors such as use-after-free and double-free bugs (Nagarakatte *et al.*, 2010; Serebryany *et al.*, 2012), and memory leaks caused by missing deallocation (Hastings and Joyce, 1992; Cherem *et al.*, 2007; Sui *et al.*, 2012; Fan *et al.*, 2019). Such bugs can lead to data corruption, information disclosure, or arbitrary code execution, and are therefore frequently exploited in real-world attacks. As a result, memory errors remain a persistent attack surface in real-world systems (Szekeres *et al.*, 2013).

From an auditing perspective, memory bugs require reasoning about

the validity of memory locations and the legality of operations performed on them. An auditor must determine, for instance, whether an object has been properly allocated and initialized, which pointers or references may refer to that object, and whether the object remains live at each use site. Accordingly, auditing for memory bugs requires collecting low-level primitive facts about pointer-related program behavior, including *points-to facts*, *aliasing facts*, and the locations of allocation, deallocation, and dereference operations. Section 4.1.1 discusses how LLM-driven pointer analysis can infer such pointer-related facts and thereby provide a foundation for memory-bug detection.

Out-of-bounds accesses are a particularly important subclass of memory bugs because they blur the boundary between memory reasoning and numeric reasoning. Although they manifest as memory bugs, their root causes often lie in incorrect index computations, erroneous pointer arithmetic, or inconsistent bounds checks. Consequently, evidence for an out-of-bounds access must connect two classes of facts, namely facts about the target memory object and its size, and numeric facts about indices and offsets. Importantly, out-of-bounds accesses are not a uniform class of local access errors, but rather a broad and heterogeneous family of bugs spanning both the spatial and temporal dimensions of memory safety. For such challenging memory bugs, auditors typically need to analyze pointers and numeric values precisely, because the two domains are often deeply intertwined. Section 4.1.4 presents the details of LLM-driven invariant generation, which can derive expressive program invariants over values from multiple domains, including both pointer and numeric domains.

2.1.3 Taint-Style Bugs

Taint-style bugs arise when data crosses implicit trust boundaries without adequate validation, sanitization, or authorization checks. In the common case, untrusted or adversarial data propagates from external inputs to security-sensitive operations. Conversely, confidential data may also flow to publicly observable outputs without proper protection. Such data, commonly referred to as *tainted values*, may originate from user inputs, network messages, file systems, or inter-process commu-

nication channels. Typical instances include command injection (Su and Wassermann, 2006), cross-site scripting (Su *et al.*, 2023), and SQL injection (Gowtham and Pramod, 2021). These bugs allow attackers to influence control flow, data access, or resource usage in unintended ways, often leading to privilege escalation, information disclosure, or malicious command execution. Unlike memory bugs, taint-style bugs usually do not cause immediate runtime failures. Instead, they compromise security guarantees by violating implicit trust boundaries and input validation assumptions embedded in the program.

Statically detecting taint-style bugs is typically formulated as a data-flow problem rather than as a classification of tainted values. An auditor must determine how tainted values are preserved, transformed, or terminated across possible executions, and whether the resulting flow violates the program’s intended trust boundaries. In real-world codebases, this task requires three supporting analyses.

First, to handle real-world programs, tainted values often need to propagate through mutable heap state and indirect memory accesses. For example, tainted values may be written into a field, array element, or heap object through one reference and later read through another aliased reference. Such flows are implicit with respect to syntactic variable names, and a def-use analysis that ignores aliasing may fail to connect the write and the read. Therefore, a precise taint analysis often relies on *pointer analysis* to identify points-to and aliasing relationships relevant to taint propagation. Missed aliases lead to false negatives, while spurious aliases inflate the set of tainted values. Section 4.1.1 discusses how neurosymbolic analysis can reason about pointer properties, including aliasing.

Second, taint-style bugs in real applications commonly depend on complex control flow. The propagation of a tainted value may be distributed across different functions, modules, framework callbacks, and library boundaries. Moreover, it can be governed not only by direct function calls but also by virtual dispatch, asynchronous callbacks, event-driven handlers, and other implicit control-transfer mechanisms. In languages where such bugs are prevalent, such as Java and JavaScript, these features obscure which code may execute and in what order, and hence which paths may carry tainted values. Therefore, taint analysis

requires precise interprocedural control-flow analysis. Call graph analysis is a central component of this reasoning because it approximates the caller-callee structure over which tainted values are propagated. In object-oriented programs, precise call graph construction further depends on points-to analysis to resolve dynamic dispatch. Imprecise control-flow or call graph analysis introduces spurious paths or caller-callee edges that cause tainted values to propagate along infeasible paths, inflating both the analysis cost and the false-positive rate. Section 4.1.3 surveys LLM-aided call graph analysis, a key instance of control-flow reasoning for improving the precision of taint-style bug detection in such settings.

Third, even when a candidate flow path is recovered, determining whether it constitutes a bug depends on the security meaning of the APIs and operations along that path. Auditors must know which program inputs should be treated as sources of untrusted data, which operations are security-sensitive sinks, and which checks or transformations neutralize taint. Therefore, accurate *taint specifications* are indispensable (Chibotaru *et al.*, 2019). Among these specifications, sanitizer identification is especially challenging. Unlike sources and sinks, which often follow recognizable API patterns, sanitization logic is typically application-specific. The absence of comprehensive sanitization specifications is a major contributor to false positives in practice, as legitimate flows that pass through custom validation are incorrectly flagged. Section 4.2.1 presents neurosymbolic techniques for inferring such specifications from documentation and other non-code artifacts.

In this sense, taint-style bugs are closely related to certain memory bugs, such as null pointer dereferences and use-after-free errors, because they all require reasoning about specific forms of data flow. In general, detecting such bugs benefits from pointer analysis, control-flow analysis, and specification inference, but the main difficulties differ because the relevant program features and bug semantics differ. For example, in memory-bug detection, specification inference often focuses on identifying memory-management functions rather than inferring sanitization specifications. Similarly, in C/C++ programs where memory leaks and other memory-management errors are common, function pointers and the induced indirect calls are often the central targets of call graph

analysis. Thus, the same primitive analyses support both taint-style and memory-bug auditing, while their concrete bottlenecks and specification targets are determined by the bug class under consideration.

2.1.4 Functional Bugs

Unlike the three bug classes discussed above, functional bugs are not primarily violations of low-level safety or security properties. Instead, they arise when an implementation deviates from its intended semantics, i.e., the functional-correctness specification of what a program, component, or service is expected to compute or guarantee. Such intent can be expressed explicitly at multiple abstraction levels. At the local program level, functions, methods, and classes may be specified using assertions, pre-/post-conditions, invariants, behavioral contracts (Meyer, 1997), and refinement or liquid types (Rondon *et al.*, 2008). At the application and service level, intent may be captured by domain rules, data-integrity constraints (Huang *et al.*, 2023a), and performance or latency expectations (Yang *et al.*, 2018). Such intent may also remain implicit in design assumptions and business logic.

A functional bug is a violation of these intended properties. At lower abstraction levels, functional bugs include incorrect return values, violated object invariants, missing precondition checks, or state updates that fail to establish required postconditions. At higher abstraction levels, they manifest as violations of domain or business logic, such as smart-contract logic vulnerabilities (Sun *et al.*, 2024), missing database integrity constraints (Huang *et al.*, 2023a), performance anti-patterns that violate expected response-time behavior (Yang *et al.*, 2018), and configuration or coordination mistakes that degrade service availability. In practice, functional bugs often manifest as service-level failures rather than localized faults. For instance, large-scale outages can be triggered by subtle errors in DNS resolution, configuration logic, or coordination protocols. The programs involved may remain memory-safe and type-correct, yet still violate higher-level assumptions about availability and system behavior. These cases illustrate that functional bugs can have severe real-world impact despite leaving classical safety properties intact.

This taxonomy should not be interpreted as a strict partition of bug

types. Rather, it highlights the dominant semantic property that makes each class difficult to audit. A bug may simultaneously involve low-level safety, security, and functional correctness concerns. For example, several numeric bugs, such as numerical instability in machine learning systems or scientific computing programs, may ultimately manifest as incorrect output values and hence as failures to implement the intended functionality. Similarly, a memory or taint-style bug may become functional when its externally visible consequence is a violation of service behavior, data integrity, or application-specific requirements. Functional bugs therefore provide a broad category for reasoning about intent violations, while still overlapping with the other bug classes when their effects are interpreted at the level of program functionality.

Detecting functional bugs requires reasoning over a broad and heterogeneous set of program properties, often spanning data flows within code, configuration files, and interactions with external systems. While violations of well-specified assertions or contracts can, in principle, be checked using conventional verification techniques, many functional bugs arise precisely because such specifications are incomplete, implicit, or informal. Consequently, *specification inference* becomes a central prerequisite. Auditors must recover the intended behavior, invariants, environmental assumptions, and domain-specific constraints against which the implementation should be judged. Such specifications are often scattered across design documents, comments, and commit messages, requiring auditors to reason over natural-language artifacts in addition to code. We discuss neurosymbolic techniques for extracting and structuring such specifications from non-code artifacts in Section 4.2. As a result, functional bug detection is challenging not because the resulting properties are inherently unverifiable, but because it demands integrated reasoning over diverse and often informal sources of intent, making it one of the most complex and least standardized areas of automated code auditing.

2.2 Static Analysis and Its Challenges

This section discusses several representative static analysis techniques that have been applied to code auditing in real-world software systems.

We organize the discussion around four representative automated static analysis techniques, namely symbolic execution, abstract interpretation, data-flow analysis, and pattern-matching-based analysis. For each category, we present an overview of its methods, summarize representative studies, and highlight the classes of bugs for which it is most effective, following the taxonomy introduced in Section 2.1.

2.2.1 Static Analysis Techniques

Symbolic execution is a static analysis technique that systematically explores execution paths by treating program inputs as symbolic variables rather than concrete values (King, 1976; Baldoni *et al.*, 2018; Boyer *et al.*, 1975; Howden, 1977). Along each explored path, the analysis accumulates branch conditions and encodes the semantics of program statements as logical constraints, thereby constructing symbolic states that represent program memory under these constraints. To assess whether a bug may occur at a particular program location, the triggering condition of the bug is conjoined with the corresponding symbolic state. Symbolic execution is well-suited for detecting bugs that manifest under specific input conditions, such as null pointer dereferences, buffer overflows, and divide-by-zero errors, by systematically exploring path conditions and solving constraints to generate triggering inputs. Tools such as KLEE (Cadar *et al.*, 2008) and Symbolic PathFinder (Havelund and Pressburger, 2000) have demonstrated effectiveness in automatically uncovering deep semantic errors in real-world software systems. Despite its precision and capacity to reason about complex execution paths, symbolic execution is fundamentally limited by the path explosion problem, which severely restricts its scalability. To address this challenge, a variety of techniques have been developed, including heuristic search strategies (He *et al.*, 2021), state merging (Kuznetsov *et al.*, 2012; Torlak and Bodík, 2014; Porncharoenwase *et al.*, 2022), compositional symbolic execution (Löw *et al.*, 2025), parallelism (Wei *et al.*, 2023; Bucur *et al.*, 2011), and compilation (Wei *et al.*, 2020; Wei *et al.*, 2023; Poeplau and Francillon, 2020). Typically, symbolic execution can explore only partial program paths, especially in the presence of loops and recursive calls, yielding an *under-approximated* view of program behavior.

Abstract interpretation provides a principled framework for soundly approximating program semantics through the design of abstract domains and abstract transformers over program operations (Cousot and Cousot, 1977a; Cousot and Cousot, 2014). In contrast to symbolic execution, which reasons about a subset of feasible paths, abstract interpretation computes abstract states that collectively over-approximate the set of all reachable concrete states. This over-approximation enables analyzers to establish the absence of certain classes of bugs when no error state is reachable in the abstract semantics. It has been applied to a broad range of bug classes, including numeric errors such as divide-by-zero, memory bugs such as null pointer dereferences and buffer overflows, and selected forms of resource mismanagement such as file-descriptor leaks. Industrial analyzers such as Astrée demonstrate the practical impact of abstract interpretation (Cousot *et al.*, 2005), especially in safety-critical settings where soundness guarantees are central. The main challenge lies in balancing precision, scalability, and expressiveness. Coarse abstractions scale well but may introduce many spurious alarms, whereas precise abstractions can be computationally expensive and difficult to apply to large programs. Consequently, practical abstract interpreters often rely on domain-specific abstractions, widening/narrowing strategies, and refinement mechanisms tailored to particular properties or program classes. Techniques such as demand-driven abstraction refinement further improve this balance by selectively increasing precision only where it is needed to discharge alarms or prove target properties (Stein *et al.*, 2019).

Data-flow analysis (Kildall, 1973) comprises a broad family of program analysis techniques for reasoning about how values and program states propagate across instructions and procedures (Reps *et al.*, 1995; Sagiv *et al.*, 1996; Rountev *et al.*, 2008; Späth *et al.*, 2017). A data-flow analysis typically maintains a set of program facts and updates them according to the semantics of each instruction, often expressed through Gen/Kill functions that specify which facts are produced or invalidated. Classical instances include reaching definitions, constant propagation, and live-variable analysis, which have long served as core components of compiler optimization. Beyond these foundational analyses, specialized

forms of data-flow reasoning have been widely used for security and correctness checking. Taint analysis (Arzt *et al.*, 2014) tracks whether values derived from untrusted sources can reach security-sensitive sinks, thereby identifying potential injection vulnerabilities and various other forms of taint-style bugs. Typestate analysis (Fink *et al.*, 2006) instead reasons about whether objects are used according to valid operation protocols, such as ensuring that files are closed after being opened or that allocated resources are eventually released. A central practical challenge in precise data-flow analysis is pointer reasoning. Because data-flow facts often depend on aliasing relationships, many analyzers adopt layered designs in which points-to or alias information is computed before the target analysis (Sui and Xue, 2016). However, this design can lead to the so-called “pointer trap” (Shi *et al.*, 2018), i.e., improving pointer precision can substantially improve downstream data-flow accuracy, but the pointer analysis itself may dominate the overall cost. This creates a persistent tension between precision and scalability. To mitigate this tension, recent techniques employ demand-driven analysis, which computes aliasing or data-flow facts only when they are needed (Späth *et al.*, 2016), or exploit parallel and incremental design to scale data-flow analysis to large programs (Liu *et al.*, 2019).

Pattern-matching-based analysis detects bugs by searching specific syntactic or semantic patterns upon program representations, such as abstract syntax trees (ASTs) and control-flow graphs (CFGs). Tools such as PMD (PMD Project, 2026) and FindBugs/SpotBugs (SpotBugs Community, 2026) exemplify this family of techniques. Specifically, PMD applies extensible rules over source-level ASTs, whereas FindBugs and its successor SpotBugs identify recurring bug patterns in Java bytecode. Mainstream integrated development environments (IDEs), including JetBrains products, also provide structural-search facilities and inspection rules that can capture common coding idioms and recurring bugs. Compared with semantics-based analyses such as symbolic execution, abstract interpretation, and data-flow analysis, pattern-matching-based techniques are tailored to specific syntactic patterns, yet are lightweight, efficient, and easy to deploy. In principle, rules can be designed to target concrete patterns from any of the bug classes discussed above,

including numeric, memory, taint-style, and functional bugs. In practice, however, such rules often cover highly specific cases, and their effectiveness is bounded by the coverage, expressiveness, and quality of the predefined templates. Pure pattern matching is especially limited because it generally lacks deep interprocedural semantic reasoning about aliasing, path feasibility, and data dependencies. Consequently, pattern-matching-based approaches may produce substantial false positives when pattern matches do not correspond to semantically feasible bugs. They are therefore best viewed as practical complements to more principled program analyses, since they can efficiently detect common and recurring bugs but are less suitable for discovering bugs whose manifestation depends on complex semantic context.

Summary. The four categories of static analysis techniques discussed above collectively constitute the methodological foundation of automated code auditing. Each technique entails a distinct balance among precision, recall, efficiency, and scalability. Symbolic execution offers fine-grained, per-path reasoning and high precision, enabling the detection of deep semantic errors, although it generally under-approximates program behaviors due to the path explosion problem. Abstract interpretation, in contrast, provides sound over-approximation guarantees and has demonstrated industrial success in safety-critical domains, albeit at the cost of precision-efficiency trade-offs. Data-flow analysis occupies a versatile middle ground, supporting both classical compiler optimizations and advanced security analyses. Its effectiveness, however, is tightly coupled with the precision of pointer analysis, which often introduces significant computational overhead. Pattern-matching-based analysis, while lightweight and efficient, is limited to recurring, rule-specified patterns and therefore serves as a complement rather than a substitute for more principled analyses. Additionally, other classes of program analysis techniques, such as fuzzing and dynamic testing, fall outside the scope of this survey. These approaches rely on program execution to explore runtime behaviors and therefore differ fundamentally from our focus on practical static analysis methods that operate without program execution and are designed for scalable bug detection across real-world software systems.

Table 2.1 lists a representative set of static analysis tools, categorized by their underlying techniques and supported programming languages. In practice, industrial analyzers are rarely based on a single technique. For example, Coverity integrates abstract interpretation with symbolic execution and data-flow analysis for bug detection (Bessey *et al.*, 2010), illustrating that real-world static analyzers are fundamentally hybrid systems rather than instantiations of any single theoretical framework.

2.2.2 Common Challenges of Static Analysis

While static analysis has achieved significant progress, existing techniques continue to encounter fundamental limitations that hinder their effectiveness and widespread adoption. Well-known challenges, such as limited precision, poor scalability, and substantial analysis overhead, have been extensively studied in both academic research and industrial practice. Beyond these established concerns, however, several overlooked yet equally critical barriers remain unsolved. In particular, we emphasize three underappreciated obstacles: (1) intermediate representation reliance, (2) limited customizability, and (3) semantic barrier. These challenges profoundly impact the applicability and usability of static analysis tools in real-world software engineering environments.

Intermediate Representation (IR) Reliance. As demonstrated in Table 2.1, most static analysis techniques operate on a specific intermediate representations, such as the Jimple IR for Java programs (Bodden, 2012) or the LLVM IR for C/C++ programs (LLVM Project, 2026). Specifically, IRs provide normalized and structured representations of programs. They are typically used as internal data structures by compilers, eliminating unimportant syntactic details while preserving essential semantics. By operating on uniform IRs rather than unstructured textual source code, analysis can abstract away language-specific syntax and focuses on semantic properties relevant to bug detection. For instance, Soot performs analyses over the Jimple IR (Bodden, 2012), and Meta Infer intercepts compilation commands to extract the IR required for its checks (Facebook, 2022).

However, the requirement for IRs also introduces fragility into large-

Tool	Method	Language	IR Reliance	Customizability
KLEE	SymEx	C, C++	Yes	No
Symbolic PathFinder	SymEx	Java	Yes	No
Clang Static Analyzer	SymEx	C, C++	Partial	Partial
Astrée	AbsInt	C, C++	Yes	No
Infer	AbsInt	Multiple	Yes	No
Frama-C/Eva	AbsInt	C	Partial	No
FlowDroid	DFA	Java	Yes	No
SVF	DFA	C, C++	No	No
Pinpoint	DFA	C, C++	Yes	No
FindBugs/SpotBugs	PM	Multiple	Partial	Yes
PMD	PM	Multiple	No	Yes
Semgrep	DFA, PM	Multiple	No	Yes
CodeGuru	DFA, PM	Java, Python	Partial	No
CodeQL	DFA, PM	Multiple	Partial	Partial
Coverity	Hybrid	Multiple	Partial	Partial

Table 2.1: A comparison of representative static analysis tools. SymEx, AbsInt, DFA, and PM stand for symbolic execution, abstract interpretation, data-flow analysis, and pattern-matching-based analysis, respectively. *IR Reliance* indicates whether the tool relies on IRs beyond the AST. *Customizability* denotes the extent to which developers can define detectors for project-specific bugs without substantially modifying the analyzer. “Yes” indicates user-level customizability. “Partial” indicates customizability that requires substantial expertise. “No” indicates that the tool primarily targets fixed families of bugs.

scale and evolving software ecosystems. First, most IRs, except for ASTs, cannot be easily obtained during early stages of development, causing analyzers to miss opportunities to surface bugs in earlier, incomplete versions. Second, this requirement complicates analysis across diverse build systems and compilers. More critically, many analyzers are compatible only with specific IR formats, making them highly sensitive to compiler versions and build systems. When projects adopt newer compilers that emit updated IR formats, as in recent Linux kernel releases, existing tools may fail due to incompatibility. This phenomenon, often referred to as the IR version trap (Zhang *et al.*, 2024), poses a significant obstacle to the longevity and portability of static analysis tools. Finally, the rigid dependency on IRs also impedes the auditing of rapidly evolving or automatically generated code, such as AI-suggested code fragments. Although recent research has begun to explore strategies for loosening

this dependency, for example by reconstructing partial or incomplete programs into analyzable forms, mainstream analyzers remain brittle, with their effectiveness tightly coupled to the successful generation of the required IR.

Limited Customizability. Another major limitation of static analysis is its limited customizability. Existing analyzers often lack mechanisms that allow users to adapt analyses to new bug classes or project-specific semantics without substantial engineering effort. As discussed in Section 2.1, different bug classes impose distinct reasoning demands and require different forms of evidence. Real-world projects further specialize these demands through programming languages, frameworks, APIs, and coding practices, making project-specific customization essential. For example, the MISRA C++ standard enumerates a broad spectrum of bugs ranging from undefined behavior in pointer arithmetic to subtle resource management errors (MISRA, 2008), highlighting the diversity of bug patterns even within a single language. Yet most analyzers are built around fixed rule sets or narrowly scoped property families. FlowDroid, for instance, is primarily designed for taint-style bugs (Arzt *et al.*, 2014), whereas Meta Infer targets memory errors such as buffer overflows (Facebook, 2022).

Extending analyzers to new bug patterns typically requires custom rules or plugins, which in turn demand expertise in compiler infrastructures, IRs, and domain-specific rule languages. As a result, developers must either design new analyzers from scratch or substantially modify existing ones. Empirical evidence shows that more than 70% of developers who have used static analyzers express dissatisfaction with the limited support for desired project-specific adaptations (Johnson *et al.*, 2013). Although recent tools such as CodeQL provide query-based customization (Avgustinov *et al.*, 2016), developers must still master a domain-specific query language (CODEQL, 2024) and a large API ecosystem, thereby introducing a steep learning curve. Other tools, such as FlowDroid, SVF, and Pinpoint, allow developers to customize data-flow specifications, while such customization mainly varies parameters within an existing analysis rather than enabling developers to define new semantic properties across diverse bug classes.

Semantic Barrier. Another fundamental limitation of existing static analysis techniques lies in their requirement for explicit semantic modeling, often in the form of formal and well-defined semantics. Specifically, most analyzers depend on two prerequisites. The first is a formal model of bug-triggering conditions, which specifies the logical specification under which a bug may manifest, such as an arithmetic operation exceeding representable bounds, a dereference reaching a null or freed object, an untrusted value reaching a sensitive sink, or an implementation violating an intended functional property. The second is a formal semantics of program behavior, which defines the meaning (e.g., operational semantics) of program constructs, such as statements, instructions, or even APIs and libraries. The well-defined semantics indeed allow analyzers to collect and check evidence, including numeric ranges, aliasing and lifetime facts, data-flow paths, and inferred contracts. Collectively, these two components determine the semantic reasoning capacity of static analysis techniques.

In practice, however, both requirements can be difficult to fulfill. Many bug-triggering conditions are complex to formalize even for human experts, especially for functional bugs whose correctness criteria are informal or domain-specific (Zheng *et al.*, 2025a; Guo *et al.*, 2025a). Semantics integrity is also impeded by incomplete or unavailable code: A common example is the absence of third-party library implementations, which prevents tools from automatically modeling the semantics of external function calls. Although manually written library specifications are a widespread workaround, this process is labor-intensive, error-prone, and requires substantial domain expertise. Together, these limitations constrain the semantic understanding that existing auditing techniques can achieve, leaving many program bugs, particularly those involving complex logic or hidden dependencies, beyond their reach.

2.3 LLM-driven Code Analysis and Its Challenges

Recent advances in LLMs have reshaped how code auditing tasks can be approached. Rather than relying exclusively on explicitly engineered abstractions, LLM-driven code analysis explores the use of neural models to infer semantic properties from source code, documentation, and other

software artifacts. This capability is particularly relevant to the limited-customizability and semantic-barrier challenges discussed above, because many auditing tasks require project-specific evidence and informal intent that are difficult to encode manually. In this section, we first review LLM-driven code analysis techniques and then discuss the common challenges that limit their accuracy and scalability in practice.

2.3.1 LLM-driven Code Analysis Techniques

Early efforts toward LLM-driven code analysis focused on incorporating program structure into neural representations through pretraining. Models such as GraphCodeBERT extend language-model pretraining with explicit representations of program graphs (Guo *et al.*, 2020), enabling them to capture syntactic structure as well as semantic relations such as data-flow and control-flow edges among program elements (Chen *et al.*, 2024). By exposing models to these structured signals during pretraining, such approaches demonstrate that neural models can learn to predict nontrivial program properties. In parallel, reinforcement learning has been explored to guide models toward capturing additional semantic attributes, further strengthening the connection between learned representations and program semantics.

Building on these foundations, more recent work has shifted toward leveraging general-purpose, decoder-only LLMs for program analysis via prompting (Pei *et al.*, 2023; Wang *et al.*, 2024c; Cheng *et al.*, 2025). Instead of post-tuning models with task-specific supervision, these approaches rely on carefully designed prompts to elicit analytical behavior at inference time. Notably, chain-of-thought prompting enables LLMs to articulate intermediate reasoning steps, allowing them to simulate aspects of program execution, reason about control flow, and track variable states across code fragments (Fang *et al.*, 2024a). Through such prompting strategies, LLMs have been applied directly to source code for bug detection.

Compared with classical static analysis, LLM-driven code analysis offers several practical advantages. First, it does not require intermediate representations, making it applicable to evolving codebases during development. Second, analysis logic can be customized through natural-

language prompts, significantly lowering the barrier to adapting analyses to new bug patterns or property domains. Third, because LLMs are pretrained on large and diverse code and text corpora, they possess rich prior knowledge about programming idioms, common libraries, and API usage, enabling them to reason across the semantic barrier that is difficult for static analyzers to cross, such as calls to external libraries or undocumented framework behavior. These strengths make LLM-driven analysis a compelling complement to static analysis, while also raising important questions about its limitations, which we discuss next.

2.3.2 Common Challenges of LLM-driven Code Analysis

Despite these encouraging capabilities, LLMs alone are far from sufficient for effective code auditing. Their underlying neural architecture, training, and inference paradigm are fundamentally misaligned with several core requirements of code auditing. In what follows, we illustrate three representative failure modes that systematically limit the effectiveness of LLM-based code analysis in realistic settings.

Sensitivity to Nonlinear Program Structure. Unlike textual natural language, program behaviors are dynamic and induced by nonlinear/nonlocal control-flow and data-flow structures. Auditing tasks such as detecting resource leaks, null pointer dereferences, or use-after-free errors require reasoning about the reachability of specific program points and the propagation of values along particular execution paths. However, LLMs process code primarily as linear token sequences and do not explicitly construct or reason over CFGs or data-flow graphs (DFGs). As a result, they are prone to implicitly assuming “typical” execution patterns from the pretraining dataset rather than faithfully tracking feasible paths.

This limitation is illustrated by the example in Figure 2.1, where a `goto`-based error-handling structure causes execution to bypass resource deallocation. Correctly identifying the resulting memory leak requires recognizing that the jump skips the `free` statement on that particular path. In contrast, an LLM may hallucinate a more benign control flow in which the deallocation is executed, simply because such patterns are

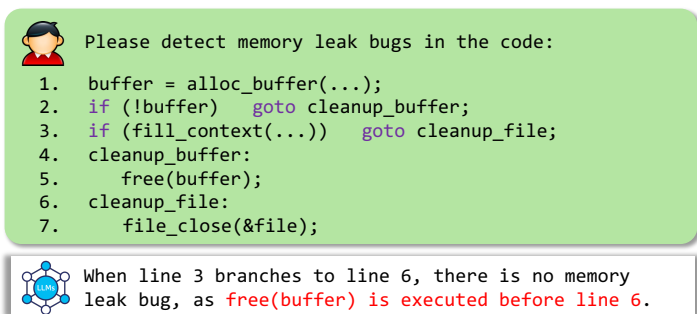


Figure 2.1: A false negative when analyzing multiple if and goto statements

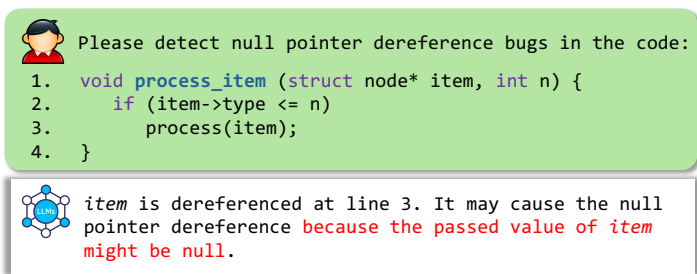


Figure 2.2: A false positive caused by the lack of calling context

common in similar code. This tendency leads to false negatives. Real bugs are missed because the model implicitly repairs or simplifies the control-flow structure in its internal reasoning.

Dependence on Missing Calling Contexts. As stated in Section 2.1, auditing complex properties requires interprocedural reasoning of the program. Whether a pointer can be NULL, whether a buffer index is in bounds, or whether a resource has been initialized often depends on how a function is called, the values of global variables, and protocol-level invariants that are not visible within a single function body. Human auditors routinely consult call sites, documentation, and usage patterns to understand such preconditions.

In contrast, LLMs are typically prompted with isolated code fragments. When crucial calling contexts are absent, the model cannot distinguish between values that are unconstrained and values that are

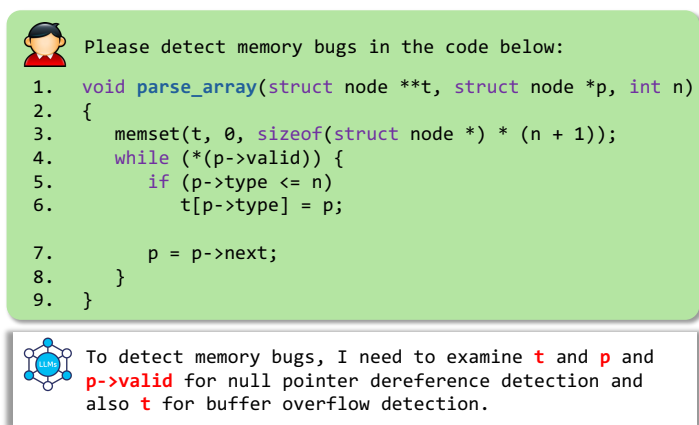


Figure 2.3: The incomplete reasoning caused by insufficient test-time compute

constrained by external guarantees. As shown in Figure 2.2, the pointer `item` is dereferenced inside `process_item`, but its nullability is determined entirely by its callers. Without access to these callers, an LLM tends to treat `item` as potentially null and will often report a null pointer dereference even when none can occur in any feasible execution. Such context-insensitive reasoning systematically produces false positives, undermining the practical usefulness of LLM-based analysis.

Insufficient Test-time Compute for Multi-fact Reasoning. Realistic auditing tasks rarely involve a single, isolated risk. A single function may simultaneously manipulate multiple pointers, indices, and buffers, each governed by different constraints that interact in subtle ways. Accurately reasoning about such code requires maintaining and updating several interdependent invariants, such as the relationships among array sizes, loop bounds, pointer validity, and data structure topology.

Although modern LLMs support extended context windows, their test-time compute remains limited. When faced with code such as the example in Figure 2.3, the model must track the values of `t`, `p`, `p->valid`, and `p->type`, while simultaneously reasoning about array bounds and pointer dereferences. In practice, the model tends to allocate its limited reasoning budget unevenly, often focusing on a subset of potential issues

while neglecting others. Consequently, some real bugs are overlooked, while other potential risks are conservatively over-approximated. This imbalance leads to both false negatives and false positives, and cannot be reliably remedied by simply increasing the length of prompts or the number of reasoning steps.

2.4 Summary

This chapter revisits code auditing from two complementary perspectives, namely static analysis and LLM-driven code analysis. Static analysis provides precise, structured, and formal reasoning over program semantics. However, its effectiveness is often constrained by several limitations, such as IR reliance, limited customizability, and the semantic barrier. In contrast, LLM-driven code analysis offers flexibility and adaptability. It can operate without IRs, be customized through natural language descriptions, and leverage rich prior knowledge to reason across incomplete code, external libraries, and non-code artifacts. Unfortunately, its reasoning inherently suffers from instability, hallucination, and the lack of formal guarantees.

These observations indicate that classical static analysis and LLM-driven code analysis should not be regarded as competing alternatives, but rather as complementary techniques whose relative advantages depend on the auditing context. For example, in early-stage development scenarios, where code may be incomplete or difficult to build, LLM-driven analysis can provide valuable semantic insights, while lightweight static analyses can supply structural grounding. Moreover, the ability of LLMs to perform semantic reasoning over multimodal software artifacts can help static analysis overcome the semantic barrier, for instance by interpreting library APIs and system-level specifications. Conversely, when the compilation pipeline is available and stable, static analyses perform deeper semantic reasoning over IRs, reducing the reliance on LLM-based inference for low-level properties such as control flow, aliasing, and numeric properties. Effective code auditing therefore requires selectively combining these techniques based on the constraints and objectives of the analysis scenario.

From a broader perspective, LLM-driven code analysis is closer to

how human auditors inspect code. It leverages an understanding of programming-language semantics and natural-language intent to reason about a restricted scope of code. Static analysis, in turn, provides the auxiliary tooling that human auditors rely on to systematically search for critical program constructs, derive relevant semantic properties, and validate or refute hypotheses with structured evidence. This analogy naturally motivates an *agentic* formulation of code auditing, in which an auditing agent orchestrates LLM-based reasoning and static analyses to emulate human auditing workflows across the software development lifecycle. Here, the agentic formulation does not presuppose full autonomy or a planning-based architecture. Instead, it provides a unifying abstraction for the diverse forms of orchestration observed in practice, including hand-designed fixed workflows and autonomous LLM agents. What makes these designs neurosymbolic in this survey is not merely the coexistence of an LLM and a static analyzer, but the multi-step coordination of heterogeneous tools, artifacts, and memory in an evidence-seeking workflow. Such an agent-centric view provides a unifying framework for understanding and designing diverse instances of neurosymbolic code auditing, which we develop in detail in [Chapter 3](#).

3

Code Auditing From an Agentic Perspective

By examining how human auditors conduct manual code auditing, we observe that auditing proceeds as a multi-step and often iterative reasoning process rather than a single-step reasoning process. Specifically, human auditors analyze multiple bug-relevant program properties and leverage these properties as evidence to determine the presence of bugs. Throughout this process, developers reason over multimodal software artifacts. In addition to source code, they routinely exploit non-code artifacts, such as library documentation, project histories, and community discussions, to infer API specifications, bug patterns, and heuristic strategies for auditing. They also rely on diverse tools, ranging from IDE-based navigation and search to existing auditing tools, to achieve effective code auditing.

Motivated by these observations, this chapter formulates an agentic perspective on code auditing. As shown in Figure 3.1, the auditor interacts with an environment of code and non-code artifacts through tool use and maintains auditing evidence in memory. This abstraction lays the foundation for Chapter 4, which introduces LLM-driven primitive analyses, and Chapter 5, which presents neurosymbolic designs that instantiate these primitives into end-to-end code auditing systems. In

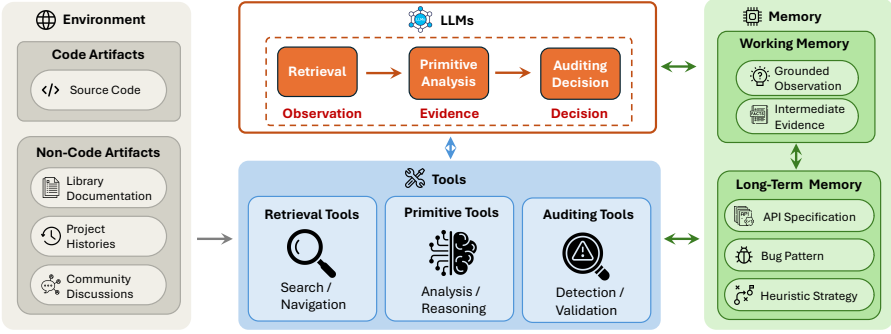


Figure 3.1: The workflow of neurosymbolic code auditing from an agentic perspective

this chapter, we introduce the characteristics of the environment, tool use, and memory for code auditing from an agentic perspective.

3.1 Environment of Code Auditing

Code auditing is inherently situated in a multimodal software environment that extends beyond source code alone. In practice, human auditors reason over a diverse range of software artifacts, which can be broadly categorized as *code artifacts* and *non-code artifacts*. Typically, human auditors directly analyze the source code of a software system during code auditing. Non-code artifacts, including library documentation, project histories (e.g., commit histories, issue reports, and patches), and community discussions (e.g., posts on Stack Overflow), provide complementary contextual information that is often indispensable for grounding auditing decisions in intended semantics and known failure patterns.

These categories of artifacts expose distinct yet complementary *primitive properties* that guide auditing decisions. Source code encodes rich semantic information about program behavior, such as data-flow facts (Bodden, 2012), typestate transitions (Fink *et al.*, 2006), and resource usage patterns (Wang *et al.*, 2023b), which together characterize how the program executes. Library documentation further refines this view by specifying API-level contracts, preconditions, and side effects that govern the correct use of external libraries. Project historical

data, such as patches and issue reports, capture recurring bug patterns, common failure modes, and the causal logic underlying past defects.

From this perspective, code auditing can be viewed as an iterative interaction between the auditor and the surrounding environment. Auditors continuously inspect, analyze, and collect primitive properties from multimodal artifacts, gradually forming evidence to support or refute the presence of bugs. The outcome of this process is not only a bug report but also an accompanying explanation linking derived program properties to violated specifications or known bug patterns. Notably, many of these properties cannot be inferred from an isolated artifact alone. Instead, they emerge only when code, documentation, usage context, and project history are examined together. The multimodal artifacts in the environment enable auditors to cross the semantic barrier imposed by any single artifact, such as source code alone, and thereby audit more effectively. To analyze different kinds of artifacts and extract these primitive properties, auditors rely on a range of analysis tools and maintain intermediate auditing state in memory across analysis steps. The roles of these tools and agentic memory are discussed in Section 3.2 and Section 3.3, respectively.

3.2 Tools for Code Auditing

To audit different forms of software artifacts effectively, human auditors rely on a range of tools to retrieve, inspect, and analyze the portions of the environment that are relevant to a bug candidate. Based on recurring behavioral patterns in manual code auditing, we categorize these tools into three classes, namely *retrieval tools*, *primitive tools*, and *auditing tools*. Specifically, retrieval tools acquire relevant context from artifacts, primitive tools derive semantic properties from retrieved contexts, and auditing tools generate bug reports at a larger scale. Together, they support observation and verification within the broader evidence-seeking workflow and accumulate the evidence needed for final bug detection decisions.

3.2.1 Retrieval Tools

In practice, developers often audit code by manually browsing and navigating a codebase across different stages of software development. As any single artifact exposes only a partial view of the program's behavior, missing context renders an incomplete view of the program's behavior. Retrieval tools are therefore essential for crossing these artifact boundaries and assembling the context required for reliable auditing decisions.

IDEs, such as IntelliJ and VS Code, provide essential navigation and querying capabilities that allow human auditors to retrieve precise program fragments on demand. Commonly used functionalities include *Definition Lookup*, which retrieves declarations of functions, types, and variables to establish their scope and intended role, *Call Graph Navigation*, which enables traversal between callers and callees to expose interprocedural dependencies, *Use-Site Analysis*, which identifies all references to a variable or function and supports reasoning about data-flow facts and object lifecycles, and *Textual and Structural Search*, which allows auditors to locate suspicious constructs using keyword-based or syntactic patterns. At a lower level, many of these capabilities are exposed through the Language Server Protocol (LSP) (Gunasinghe and Marcus, 2021), which provides a uniform interface for querying program structure and semantic information. As a result, any LSP-compliant implementation can serve as a retrieval tool, enabling agents to collect relevant program constructs and surrounding contexts in a manner analogous to IDE-assisted auditing.

Beyond local IDE support, *codebase-wide search platforms*, such as GitHub code search or Sourcegraph (Sourcegraph, Inc., 2026), enable human auditors to retrieve relevant code and documentation across large codebases. External web search plays a similar role for non-code artifacts when auditors need library documentation, issue discussions, bug reports, or usage examples that reside outside the local codebase. *Version control systems* further enrich retrieval by providing historical context. Commands such as `git blame` and `git log` allow auditors to trace the evolution of suspicious code, revealing when it was introduced, how it has changed over time, and whether it has been associated with

previous bug fixes or regressions.

Modern coding agents operate in a closely related way. Systems such as Claude Code (Anthropic, 2026a) and Codex (OpenAI, 2026a) typically do not reason over an entire codebase at once. Instead, they iteratively issue retrieval actions within the environment, including grep-like text search, file-system traversal, targeted file reads, version-control queries, and, when necessary, web retrieval for external documentation. In particular, these agents often depend on customized grep-like searches over code to quickly surface candidate definitions, call sites, and error paths. In practice, such lightweight retrieval operations are often more effective for interactive auditing than heavier static analyses or AST-based structural search, because they are easy to compose, require little setup, and adapt well to incomplete, evolving, or heterogeneous codebases.

3.2.2 Primitive Tools

Primitive tools capture recurring analytical subprocesses employed by human auditors during manual code auditing. Given the context retrieved through retrieval tools, auditors derive specific primitive program properties and incrementally store them in the memory of an auditor. This process is repeated until the accumulated properties are sufficient to determine whether the potential bug exists. We formalize this iterative analysis workflow as the use of *primitive tools*. Depending on the nature of the analyzed artifacts, we distinguish between primitive tools that operate on code artifacts and those that operate on non-code artifacts.

On code artifacts, primitive tools correspond to analyses that infer semantic properties of program behavior directly from source code. Examples include pointer analysis (Chen *et al.*, 2024), dependency analysis (Wang *et al.*, 2024c), call graph analysis (Cheng *et al.*, 2024), and invariant generation (Pei *et al.*, 2023), which support the detection of different types of bugs. For example, auditors reason about program dependencies to determine whether a pointer may be null at a dereference site and infer loop invariants to justify whether a given assertion always holds within a loop.

On non-code artifacts, primitive tools support the extraction of semantic specifications that are not explicitly encoded in the program itself. Auditors infer library API specifications from documentation, usage examples, issue reports, online discussions, and other descriptive artifacts to understand intended behavior, preconditions, and error-handling requirements. These inferred specifications are then used to interpret application code, identify semantic mismatches, and recover developer intent for bug detection.

Overall, regardless of whether they operate on code or non-code artifacts, primitive tools serve the same fundamental role. They infer primitive properties from available evidence and thereby support multi-step and even iterative code auditing workflows. The main difference lies in methodological maturity. For code artifacts, many primitive analyses have long been studied in static analysis, and substantial tool support already exists, although these techniques retain several common limitations, as discussed earlier in Section 2.2.2. For non-code artifacts, by contrast, systematic methods have historically been less mature and were often framed as traditional NLP tasks, such as semantic parsing (Zettlemoyer and Collins, 2005). With the rise of LLMs, however, primitive tools can be instantiated more effectively via neurosymbolic designs, as discussed in Chapter 4.

3.2.3 Auditing Tools

While manual code auditing is indispensable for uncovering subtle or context-dependent bugs, it is often labor-intensive and difficult to scale. In practice, human auditors therefore complement manual reasoning with automated auditing tools, most commonly static analysis tools, to identify candidate bugs across large codebases. For example, auditors routinely employ static analyzers such as Clang Static Analyzer (LLVM Project, 2026), Meta Infer (Facebook, 2022), and Coverity (Bessey *et al.*, 2010) to detect memory errors and certain logic flaws. These tools correspond to the static analysis techniques summarized in Section 2.2.1.

However, auditing tools are rarely usable out of the box. Auditors typically need to perform non-trivial setup and customization, including configuring build environments, supplying library specifications,

writing project-specific rules or queries, and even choosing appropriate parameters for performance tuning. After analysis completes, the generated reports require careful manual inspection to validate findings and discard spurious warnings. As a result, auditing tools must be integrated with retrieval and primitive tools to provide sufficient context for interpretation.

In this workflow, auditing tools function as candidate bug report generators rather than final decision makers. They direct human auditors' attention to suspicious program regions, but the final decision depends on the broader auditing workflow, which combines retrieved context, primitive analyses, and accumulated memory. We will revisit this interaction between automated auditing tools and human reasoning in Section 5.2.

3.3 Memory for Code Auditing

From an agent-centric perspective, code auditors critically depend on memory that supports multi-step evidence-seeking reasoning. We distinguish between two complementary forms of memory, namely *working memory* and *long-term memory*. Specifically, working memory stores intermediate results generated by tools during the auditing of a specific codebase, while long-term memory captures reusable knowledge that generalizes across auditing tasks, versions, projects, and domains.

3.3.1 Working Memory

Working memory maintains the intermediate results of an ongoing auditing task. These include program semantic properties derived from code and non-code artifacts as well as provisional bug hypotheses. Such artifacts are typically produced incrementally through repeated invocations of retrieval, primitive, and auditing tools, and are continuously updated as the analysis progresses.

By accumulating these intermediate results as evidence, working memory enables auditors to integrate information across different program locations and abstraction levels. This integration supports long-range reasoning that is essential for validating bugs whose evidence is

distributed across functions, files, or modules. For example, assessing a potential null pointer dereference requires retaining data-flow facts about pointer initialization and propagation across call chains, while confirming a resource leak depends on remembering allocation and deallocation events observed at different points in the program.

Working memory also plays a central role when auditors interact with static analysis tools. Bug reports generated by such tools are rarely self-contained. Interpreting them requires recalling previously derived intermediate results. Determining whether a warning is a true bug or a false positive often hinges on correlating multiple intermediate artifacts stored in working memory, such as path conditions, aliasing relationships, or prior sanitization steps.

3.3.2 Long-Term Memory

Long-term memory captures knowledge that persists across auditing problems and can be reused in different projects or contexts. This includes API specifications, domain-specific bug patterns, and established reasoning strategies for particular categories of bugs, which are typically acquired through prior audits.

Long-term memory is particularly important for compensating for missing or incomplete specifications in automated tools. When analyzers lack accurate models of library APIs or domain-specific behaviors, auditors rely on previously acquired specifications and patterns stored in long-term memory to interpret analysis results and guide further investigation. For example, recognizing misuse of a security-critical API or identifying a recurring resource-management pattern often depends on prior knowledge accumulated across multiple audits.

While working memory is task-specific, the boundary between the two forms of memory is not rigid. For example, inferred API specifications and reusable reasoning patterns may move from working memory to long-term memory for future reuse. Within a single project, previously derived program facts can also persist over time and support incremental auditing. More fundamentally, working memory extends the prior knowledge of LLMs during auditing. By persistently storing knowledge derived from past audits in long-term memory, an auditing

agent can reintroduce this information through prompts and retrieved notes, and can reorganize information acquired across different contexts and time spans. This mechanism enables reasoning beyond the semantic barrier, offering an important pathway toward more effective and scalable code auditing.

3.4 Summary

In this chapter, we present code auditing as an evidence-seeking workflow from an agentic perspective. In this view, the environment provides code and non-code artifacts, tools support retrieval, primitive analysis, and automated auditing, and memory stores grounded observations, intermediate evidence, and other facts. This view reflects a key feature of manual auditing. Auditors rarely make decisions based on a single artifact or a single reasoning step. Instead, they gather evidence, infer primitive properties, test candidate hypotheses, and relate program behavior to specifications and known bug patterns.

From an agentic perspective, we can also explain how to address the challenges of static analysis discussed in Section 2.2.2. Specifically, a rich multimodal environment helps auditors cross the semantic barrier imposed by source code alone. Tool use also extends auditing beyond classical static analysis. For example, lightweight code search and LSP-based navigation can reduce dependence on IRs and lessen the customization burden often required by static analyzers. Similarly, this perspective helps address the limitations of standalone LLM-based code analysis. Instead of asking an LLM to infer all relevant semantics from a limited context window, we can retrieve missing context from the environment, use tools to derive or verify primitive properties, and store useful evidence in working or long-term memory. These mechanisms help mitigate context limitations and reduce hallucinations by grounding model reasoning in external artifacts and tool-supported observations. In addition, iterative auditing can help overcome limited test-time compute and thus scale better to large real-world programs.

Lastly, we clarify the scope of the term *agentic* in this chapter. We do not use the term to suggest that every code auditing system must be a fully autonomous agent with open-ended planning. Rather, the agentic

perspective is a broad abstraction for auditing workflows that interact with environments, use tools, and maintain memory across multiple reasoning steps. Under this interpretation, both LLM agents with planning ability and fixed auditing workflows can be regarded as agentic systems. In Chapter 5, we discuss different neurosymbolic approaches, some of which rely on the planning ability of LLMs, while others follow fixed workflows that combine LLM-based primitive analysis with classical static analysis. For this reason, we use the term *neurosymbolic code auditing* rather than *agentic code auditing* in this survey.

4

LLM-driven Primitive Analysis

In this chapter, we explore how LLMs can be leveraged to perform a series of *primitive analyses* to extract program properties relevant to different bug types. By primitive analyses, we refer to static analyses that concern basic program properties (*e.g.*, points-to, aliasing, data dependence, control dependence, caller-callee relationships, program invariants, etc.) and serve as building blocks for a concrete auditing problem. These analyses are also often combined and integrated with other tools, ultimately enabling the detection of bugs during code auditing.

As discussed in Section 3.1, the auditing environment consists of both code artifacts (*i.e.*, source code) and non-code artifacts (*e.g.*, documents and bug reports). Therefore, our discussion of primitive analyses in this chapter is structured along these two dimensions, targeting both code and non-code artifacts, respectively.

4.1 Primitive Analysis over Code Artifacts

LLMs have increasingly been applied to primitive code analysis tasks, where they infer program properties without necessarily executing the code. We organize these analyses according to a progression of program

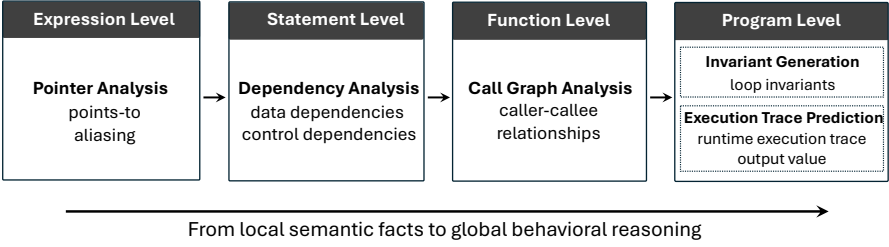


Figure 4.1: Overview of LLM-driven primitive code analysis

granularity, reflecting how relevant evidence is often accumulated during auditing. At the *expression* level, we summarize the LLM-driven analysis of pointer expressions, which discovers facts about points-to and aliasing relationships (Section 4.1.1). At the *statement* level, dependency analysis identifies data and control dependencies among statements (Section 4.1.2). At the *function* level, call graph analysis targets potential caller-callee relationships and supports interprocedural reasoning (Section 4.1.3). Finally, at the broader *program* level, invariant generation and execution trace prediction reason about properties that hold across all executions or about concrete behaviors along specific paths (Sections 4.1.4 and 4.1.5).

4.1.1 Pointer Analysis

Establishing pointer-related properties, such as points-to facts and aliasing facts, is fundamental building block in static analysis. They describe which memory locations pointers may refer to and which pointers may refer to the same memory locations. Such information supports many downstream analyses used in code auditing, including program slicing and data-flow bug detection. When inspecting pointer-heavy code, auditors do not usually compute points-to sets or alias sets explicitly, but they nevertheless maintain a mental model of which objects may be referenced and which buffers or memory regions may overlap or be disjoint with others.

In static analysis, *pointer analysis* (Smaragdakis and Balatsouras, 2015) systematically derives such pointer-related properties by computing an approximation of the memory locations that pointer expressions

may refer to. Its precision and scalability are determined by the choice of abstraction granularity, including flow sensitivity (Ye *et al.*, 2014), context sensitivity (Li *et al.*, 2020), and heap abstraction (Kanvar and Khedker, 2016). Although these analyses provide principled semantic abstractions of program statements, they require language-specific analysis infrastructure and can be difficult to apply to library functions without available implementations or to scale to programs involving sophisticated memory operations or deep calling contexts (Bastani *et al.*, 2018; Eberhardt *et al.*, 2019; Wang *et al.*, 2024a).

Recent work has therefore explored neural and LLM-assisted approaches for recovering pointer-related facts. One line of work reduces pointer reasoning to a classification problem. For example, Chen *et al.* (2024) evaluate neural code models on alias prediction. Given two pointers in source code, the model determines whether they must alias, may alias, or must not alias. Their results suggest that neural code models can capture certain local regularities of pointer semantics from real-world code, although the predicted labels do not constitute sound guarantees of aliasing facts. Similarly, LLMs have been used to infer pointer-related properties such as the nullability of C function arguments (Tullsen *et al.*, 2024). Such nullability information is useful, for example, for C-to-Rust translation, where nullable C pointers may need to be translated into Rust option types rather than ordinary references.

A complementary line of work uses LLMs not to replace static pointer analysis, but to improve the semantic modeling required by conventional analyses. For example, Cheng *et al.* (2025) use LLMs to help identify custom allocation functions in C/C++ programs. By combining static pointer analysis with LLM-based reasoning about more complex allocation patterns, their approach enriches the heap-allocation model used by a downstream pointer analysis. This illustrates a neurosymbolic use case in which the static analyzer preserves the global analysis workflow, while LLMs assist at semantic bottlenecks that are difficult to cover with hand-written rules.

Discussion. Pointer analysis is a setting where soundness matters especially strongly, because missed aliases or points-to targets can invalidate downstream reasoning about memory updates, data flow, and verifica-

tion conditions. LLM-driven pointer analysis is therefore most useful when it predicts candidate facts or enriches semantic models, rather than when it is treated as a standalone substitute for a conventional pointer analysis. In program verification, for example, LLM-derived pointer facts need to be checked or embedded inside a sound procedure before they can support proof obligations.

4.1.2 Dependency Analysis

At the statement level, code auditing often requires developers to reason about how program statements depend on one another. Given a concrete program value, auditors trace which statements produce, propagate, or use the value and which branch conditions control the reachability of these statements. *Dependency analysis* automates this reasoning by identifying data dependencies, which capture value flow from definitions to uses, and control dependencies, which capture how statement execution depends on control-flow decisions.

Static dependency analyses typically compute these relationships over dependence representations such as program dependence graphs and their interprocedural variants (Click and Paleczny, 1995; Bracevac *et al.*, 2023; Ferrante *et al.*, 1987). These representations make data and control dependencies explicit and provide a structured basis for slicing and other downstream analyses. However, they often assume that the analyzed code is complete, well-formed, and amenable to reliable discovery of def-use relations and control-flow structure (Tip, 1994; Horwitz *et al.*, 1988). These assumptions can be hard to satisfy in auditing scenarios involving partial code fragments, uncompileable code, and sophisticated language features. This motivates learning-based dependency analysis methods that approximate statement-level dependencies without requiring a fully constructed dependence graph. For example, GraphCodeBERT shows that code models can benefit from dependence-aware pretraining objectives, suggesting that models can internalize dependency information from large code corpora (Guo *et al.*, 2020).

Building on this observation, recent LLM-driven dependency analysis can be grouped into two streams, namely fine-tuned neural dependency

analysis and prompting-based dependency analysis. First, fine-tuned neural techniques use smaller code models to directly predict program dependencies, such as data/control dependency prediction (Yadavally *et al.*, 2023; Yadavally *et al.*, 2024) and sequence-to-sequence slicing with constrained outputs (He *et al.*, 2025). These methods typically formulate dependency analysis as a supervised prediction problem over statements, making them suitable for incomplete or fragment-style code where classical dependence-graph construction may fail. Second, prompting-based techniques utilize general LLMs with structured decomposition to recover data dependencies and control dependencies. These methods often combine LLM outputs with external tooling or multi-step procedures to reduce hallucination and enforce semantic constraints, as exemplified by LLMDFA (Wang *et al.*, 2024b) and NESA (Wang *et al.*, 2026). Existing empirical evidence shows that LLMs can produce approximate data/control dependencies, but the outcomes vary substantially with prompting choices and exhibit recurring error patterns (Shahandashti *et al.*, 2025). Hence, explicit verification and refinement loops are also introduced to improve the accuracy and reliability of dependency-analysis results (Chang *et al.*, 2025).

Discussion. Similar to the pointer analysis discussed in Section 4.1.1, LLM-driven dependency analysis lacks the theoretical guarantees provided by classical static analysis. Moreover, statement-level dependencies often depend on expression-level semantic facts, such as aliasing. For example, if an LLM can more accurately understand pointer-related properties within a given context, it should also be better positioned to recover data/control dependencies. Therefore, designing appropriate training tasks that help LLMs better understand memory states would be a feasible way to improve LLM-driven dependency analysis.

4.1.3 Call Graph Analysis

When analyzing a real-world codebase, human auditors routinely reason about interprocedural behavior by asking which functions may invoke which others and how execution flows across function boundaries. Call graphs provide a compact abstraction for this reasoning by

summarizing potential caller-callee relationships. To obtain all possible caller-callee pairs, *call graph analysis* needs to resolve sophisticated program constructs that may induce indirect calls, such as function pointers, higher-order functions, and dynamic dispatch. Conversely, incomplete semantic modeling may lead to missing edges and hence unsound downstream analyses.

Unlike classical static analysis, recent LLM-assisted approaches typically do not construct entire call graphs from scratch. Instead, they integrate LLM-driven semantic reasoning into conventional call graph analysis pipelines. One role is *precision-oriented pruning*. SEA (Cheng *et al.*, 2024) first obtains an initial target set for each indirect call using static analysis, and then asks an LLM to summarize the indirect-call context and candidate target functions. These summaries are then compared to filter out semantically implausible caller-callee pairs, thereby improving precision while retaining the candidate space supplied by static analysis. AutoPruner (Le-Cong *et al.*, 2022) follows a pruning-oriented view, but formulates call-graph pruning as a supervised edge classification task, i.e., given a baseline static call graph, it combines LLM-derived semantic representations of caller and callee code with graph-structural features to predict whether an edge is likely to be a false positive. In higher-order control-flow analysis, Wei *et al.* (2025) propose to use LLMs as the abstract address allocator to improve the analysis precision via context-sensitivity, meanwhile without introducing unsoundness to the analysis.

A complementary role is *recall-oriented augmentation*. CALLME targets JavaScript dynamic property-access calls (Wang *et al.*, 2025c), which are often ignored or under-approximated by classical static call graph analysis due to the difficulty of enumerating possible object and property values. Rather than merely pruning an existing graph, CALLME combines a custom static analyzer with an LLM to recover likely dynamic call edges. The language model exploits semantic signals such as variable names, natural language documentation, and calling contexts, which are difficult to encode with fixed static-analysis rules but can be informative for resolving dynamic property calls.

At the same time, empirical studies caution against treating LLMs as standalone replacements for dedicated call-graph analyzers. In a large-

scale study over Python and JavaScript, Venkatesh *et al.* (2024) report that traditional tools such as PyCG (Salis *et al.*, 2021) remain stronger at call-graph generation, while LLMs perform comparatively better on related semantic tasks such as Python type inference (Venkatesh *et al.*, 2024). This contrast is intuitive. Call graph analysis requires broad coverage of interprocedural structure and careful treatment of language and runtime semantics, whereas LLMs are strongest when a bounded context contains enough semantic, naming, and documentation cues to support local disambiguation.

Discussion. Compared with the pointer analysis and dependency analysis discussed above, call graph analysis targets a function-level program property. However, its precision often highly depends on points-to and data dependency information, including, but not limited to, aliasing relationships among function pointers and the types of specific expressions. For example, recent LLM-assisted approaches analyze function pointers to support more precise call graph analysis for the Linux kernel (Cheng *et al.*, 2024). Meanwhile, constructing an accurate call graph can also improve context retrieval, which is important for interprocedural analyses, including pointer analysis and dependency analysis.

4.1.4 Invariant Generation

When manually auditing code, human auditors reason about *program invariants*: logical properties expected to hold at designated program points across all relevant executions. Invariants operate at complementary levels of granularity. At the function level, *pre- and postconditions* characterize what a function assumes about its inputs and guarantees about its outputs, abstracting away from the concrete function body. These specifications are the cornerstone of modular code auditing, as they allow an auditor to reason about a function in isolation: a caller can be verified against the callee’s postcondition without inspecting its implementation, and the callee can be verified against its precondition without examining its call sites.

At the intra-procedural level, *loop invariants* describe properties that remain true before and after each loop iteration, allowing auditors

to abstract away from individual iterations and reason about repetitive behavior at a higher level. A sufficiently strong invariant, whether at the function or loop level, can be used to establish that required assertions continue to hold across the executions it summarizes, such as bounds checks, resource constraints, and data-structure consistency properties. When such an invariant cannot be established or does not imply the desired property, auditors are directed to the remaining proof obligations or suspicious behaviors that require closer inspection. In this sense, *invariant generation* reflects a core form of human reasoning used to explain, summarize, and validate program behavior, and it further enables the detection of a wide range of bugs related to functions and loops.

Concretely, a loop invariant is a logical assertion that remains true before and after each loop iteration. It serves as a central proof artifact in program verification (Flanagan and Leino, 2001; Cousot and Cousot, 1977b; Cousot and Cousot, 1979; Cousot and Cousot, 1977a; Floyd, 1993; Hoare, 1969; Leino, 2010). In abstract interpretation, an analysis iteratively computes an approximation of reachable states until a fixpoint is reached, with the precision of the resulting invariant depending on the chosen abstract domain and widening/narrowing strategies (Cousot and Cousot, 1977b). Additionally, automated program verification may proceed by guessing candidate invariants and checking them with a verifier. Specifically, a pool of candidate predicates is generated and then checked for inductiveness and adequacy against the verification conditions (Si *et al.*, 2018). By contrast, dynamic invariant generation methods infer likely invariants from concrete executions (e.g., from trace observations) and are often useful for documentation, test augmentation, or heuristic guidance. However, such dynamically inferred invariants do not guarantee inductiveness without additional formal verification (Ernst *et al.*, 2007).

To mitigate this bottleneck, a growing body of recent work explores LLM-assisted invariant inference as a complement to classical program verification techniques (Pei *et al.*, 2023; Chakraborty *et al.*, 2023; Liu *et al.*, 2024b). These approaches leverage the pattern-recognition and semantic generalization capabilities of LLMs to propose candidate invariants that are difficult to synthesize using purely symbolic methods.

Recent studies demonstrate that LLMs can generate plausible loop invariants from program text, execution traces, or verification conditions, and can often capture non-trivial relationships involving control-flow structure (Pei *et al.*, 2023) and memory-manipulating code (Liu *et al.*, 2024b). More recent work integrates LLM-generated invariants into verification pipelines, where candidates are subsequently checked, ranked, or refined using SMT solvers or bounded model checking, enabling a synergistic combination of neural hypothesis generation and symbolic validation (Chakraborty *et al.*, 2023). This integration reflects an emerging workflow in which LLMs assist in expanding the hypothesis space, while formal reasoning components ensure correctness.

Despite these advances, existing LLM-driven approaches typically do not address the fundamental challenges of invariant or specification inference. Generated invariants may be incomplete, overly specific, or spurious, and their effectiveness often depends on prompt design, model choice, and downstream filtering strategies. Moreover, most current techniques focus on local invariants or bounded reasoning and still struggle with codebase-wide analysis.

Discussion. Intuitively, invariant generation resembles the way human auditors reason about programs with complex control flow. For example, when analyzing a function, auditors often expect it to satisfy a specification, such as ensuring that a postcondition holds under a given precondition. To validate such a hypothesis, human auditors may first conjecture properties at specific program points and then use rigorous manual reasoning to verify them against a complete function specification. This intermediate process of conjecturing program invariants is closely aligned with current uses of LLMs for invariant generation. It is also worth noting that the properties captured by program invariants vary across applications. For example, when verifying heap-manipulating programs, the corresponding invariants may characterize complex properties over heap memory (Liu *et al.*, 2024b), including, but not limited to, points-to facts and aliasing facts studied in pointer analysis.

4.1.5 Execution Trace Prediction

Human developers often need to reason about how a program would execute under particular inputs, mentally simulating control-flow decisions and memory state updates to determine whether the program behaves as intended. However, it is not always feasible to run the program due to missing software dependencies, unsupported platforms, or prohibitive cost. *Execution trace prediction* infers dynamic program behaviors, including executed statement sequences and intermediate variable states, without actually running the program (Ni *et al.*, 2024; Xu *et al.*, 2025a). Such predictions support explaining program behavior, validating assumptions about runtime states, and identifying bugs that only manifest along specific execution paths.

Predicting execution traces is particularly challenging because runtime behavior is jointly determined by path conditions, evolving program states, and the external environment. Errors in execution trace prediction are highly non-local. A single incorrect branch decision or value update can invalidate the entire subsequent trace. This challenge is further exacerbated in real-world code by partial programs, dynamic typing, and implicit flows, where many semantic facts required for execution are difficult to discover through purely static reasoning.

Recent work explores execution trace prediction along two broad directions. The first direction injects execution semantics into the model through *trace-aware learning*. This includes approaches that condition reasoning on concrete execution traces or variable-state sequences, as well as self-training schemes that distill execution-aware rationales into the model, allowing it to internalize runtime behavior patterns beyond surface code structure (Ni *et al.*, 2024; Ding *et al.*, 2024). The second direction introduces *inference-time neurosymbolic grounding*, where symbolic structure or executable tools are used to constrain the model’s prediction. This direction can take several forms. PredEx adopts an analysis-guided decomposition strategy (Li *et al.*, 2025d), i.e., lightweight program analysis fixes deterministic execution structure whenever possible, while predictive backward slicing decomposes state prediction into smaller variable-level subproblems handled by the model. Other work uses CFG-grounded or multimodal prompting to align

reasoning with feasible control-flow paths (Le Chi *et al.*, 2025). Tool-augmented approaches further execute deterministic fragments and rely on the model only for unsupported, underspecified, or semantically ambiguous parts (Li *et al.*, 2024). Across these designs, symbolic or executable scaffolds serve as external anchors that reduce hallucinated paths and improve consistency with program semantics.

Execution trace prediction also admits several variants. For example, RepoReason evaluates whether an assertion at a specific program location holds under a given execution path (Li *et al.*, 2026). Beck *et al.* (2026) also propose the concept of a *neuro debugger*, which aims not only to use LLMs to predict execution paths and outputs under given inputs, but also to support a reverse form of execution trace prediction, i.e., inferring possible prior states or input parameters from a given current state. These techniques and variants can potentially facilitate debugging during development, especially when code fragments and environment configurations are incomplete. At the same time, evaluating LLMs on execution trace prediction provides a measure of their ability to understand code semantics and can offer useful insights for downstream tasks such as bug exploitation (Fang *et al.*, 2024b).

Discussion. Unlike invariant generation, as discussed in Section 4.1.4, execution trace prediction focuses on the runtime behavior of a program along a specific execution path under a concrete input. An accurate execution trace prediction reflects a concrete program execution. In contrast, invariant generation is concerned with whether a property holds for all inputs, or for all inputs satisfying certain conditions. Nevertheless, both techniques offer useful observations for concrete code auditing applications and are important for bug detection and diagnosis.

4.1.6 Comparisons with Classical Static Analysis

The ability of LLMs to perform primitive code analysis stems from a combination of complementary mechanisms. First, large-scale pre-training on code and code-related text exposes models to the statistical regularities of software. Second, code-specific pretraining objectives can make such semantic signals more explicit. For example, Graph-

CodeBERT incorporates data-flow and control-flow prediction (Guo *et al.*, 2020). CodeT5 introduces identifier-aware objectives to exploit the semantics conveyed by developer-assigned identifiers (Wang *et al.*, 2021). In addition, TRACED injects dynamic semantics by pretraining on execution traces (Ding *et al.*, 2024). Empirical probing studies also show that pretrained code models encode nontrivial syntactic and semantic program structures, although their semantic understanding remains uneven across tasks and models (Ma *et al.*, 2024). At the same time, LLMs can exploit informal semantic cues embedded in symbol names, API choices, comments, and coding conventions, which often reflect developer intent and engineering practice (Chen *et al.*, 2022). As a result, LLM-driven code primitive analyses are not solely based on surface-level natural-language similarity, but on a joint representation of local program structure, developer intent cues, and statistically correlated semantic patterns observed during pretraining.

When compared with classical static analyses, LLM-driven primitive code analyses exhibit fundamentally different strengths and limitations. Static analysis methods derive semantic properties by constructing explicit program representations, such as dependence graphs and abstract program states, and applying rigorous reasoning to ensure formal guarantees at a chosen abstraction level. These methods are well suited for analyses that require principled over-approximation, but they rely heavily on structured IRs, well-formalized program semantics, and carefully engineered analysis frameworks. In contrast, LLM-based approaches reformulate primitive analyses as prediction or generation problems over a restricted scope of code fragments. This design allows them to operate on incomplete, uncompileable, or fragment-level code and to approximate semantic relations that are difficult to encode with static rules. However, because their reasoning occurs in latent space, LLMs behave as black boxes without guarantees of soundness and completeness, and their errors can propagate non-locally. Consequently, LLM-driven analyses are best understood not as replacements for static analyses, but as complementary mechanisms that provide flexibility and semantic intuition where formal modeling is costly or brittle.

These observations motivate a neurosymbolic design perspective for practical code auditing tasks. Real-world auditing rarely consists

of a single monolithic analysis. Instead, it involves solving a sequence of primitive subtasks, such as identifying relevant program slices, resolving ambiguous indirect calls, or reasoning about potential runtime behaviors under partial information. In this setting, static analysis plays a role analogous to the specialized tools available to human auditors, offering structured, reusable, and verifiable results whenever formal reasoning is feasible. LLM-driven code analysis, in turn, corresponds to the primitive analysis tools formulated in Chapter 3, providing judgments and hypotheses when rules are incomplete or intent needs to be inferred. By carefully deciding how static analysis and LLM-driven analysis are combined at the level of individual primitive subtasks, these components can be composed into a coherent neurosymbolic auditing workflow. Such workflows leverage static analysis to anchor reasoning in program semantics while using LLMs to address the semantic barrier, ultimately enabling scalable and effective auditing across real-world codebases. Chapter 5 will further elaborate on this design space and its instantiations in practical code auditing systems.

4.2 Primitive Analysis over Non-Code Artifacts

As discussed in Section 3.1, effective code auditing requires reasoning not only over source code, but also over a diverse range of non-code artifacts. This section surveys primitive analyses that leverage LLMs and lightweight natural language processing techniques to derive auditing evidence from such artifacts. We organize the discussion according to the types of non-code artifacts.

4.2.1 Library Documentation

Documentation provides a major source of semantic description for code auditing, especially when analyzing code that interacts with external libraries. Many auditing tasks require knowledge of library behaviors, yet such semantics are often unavailable in a formal form or too costly to specify manually. Under this view, documentation is not merely auxiliary context for developers. Instead, it serves as a primary source of semantic description from which library specifications can be inferred.

Documentation-driven specification extraction therefore provides an important bridge between informal, developer-facing library knowledge and automated auditing of application code.

This idea predates recent LLM-based systems. Earlier NLP-based approaches already explored how to infer various forms of API specifications from natural-language documentation (Zhong *et al.*, 2009; Pandita *et al.*, 2012; Zhai *et al.*, 2016). However, these systems typically relied on task-specific linguistic patterns, semantic parsers, or hand-engineered extraction rules. Recent neural and LLM-based approaches inherit the same goal but shift the central design problem. The challenge is no longer only how to recognize particular textual patterns, but how to ground natural language understanding into stable specifications that can be consumed by downstream auditors.

Among existing studies, DAInfer is a representative example of documentation-guided specification inference designed for code auditing (Wang *et al.*, 2024a). It targets API aliasing specifications, which are crucial for library-aware data-flow analysis but are notoriously difficult to recover from library code alone. DAInfer interprets API documentation to classify memory-related operations, such as access, insertion, and deletion, and to identify the semantic entities manipulated by these APIs. On top of these abstractions, it formulates API aliasing specification inference as a neurosymbolic optimization problem, selecting a globally consistent set of aliasing facts induced by memory operations and semantic entities. This design is important because the LLM is not asked to directly emit final specifications. Instead, documentation-derived observations are converted into structured constraints, making the inferred semantics more stable and easier to integrate into existing code auditors.

Similarly, ChatDetector targets resource-management API specifications (Yang *et al.*, 2025b). It uses LLMs to recover allocation-release protocols from documentation and then constructs resource-management API constraints for misuse detection. Unlike DAInfer, ChatDetector decomposes documentation understanding into smaller and more verifiable steps and further introduces cross-validation and inconsistency checking to mitigate hallucination before the recovered constraints are used by an auditor. Compared with DAInfer (Wang *et al.*, 2024a), ChatDetector

illustrates a complementary design philosophy in which reliability is enforced through staged reasoning and validation of intermediate results rather than through a global semantic objective.

Beyond specifications for library APIs, another line of related work studies whether LLMs can extract more general formal specifications from general software documentation. Although these specifications are not always library API contracts, they expose a similar challenge that informal documentation must be converted into precise, checkable properties before it can support automated testing or auditing. For example, Li *et al.* (2025b) empirically studied the use of LLMs to translate natural-language documents into formal specifications for test automation. They found that out-of-the-box LLMs can often identify useful rules but also tend to oversimplify conditions and fabricate unsupported requirements. To mitigate these problems, they proposed a two-stage annotation-then-conversion approach, in which the model first highlights relevant facts in the text and then translates them into formal statements. This result illustrates both the promise and the current limitations of directly inferring precise formal specifications from documentation.

Discussion. Across these documentation-guided specification inference techniques, the most significant differences lie not merely in the specific specification targets, but in how language-derived outputs are grounded. DAInfer (Wang *et al.*, 2024a) relies on global consistency enforced by optimization, whereas ChatDetector (Yang *et al.*, 2025b) relies on procedural grounding through task decomposition, cross-validation, and intermediate consistency checks. The formal-specification extraction setting further suggests that even when LLMs can identify relevant semantic facts, direct end-to-end translation remains fragile without explicit grounding steps. Thus, the central design question is not whether LLMs can read documentation, but how their linguistic judgments are constrained, validated, and translated into stable artifacts consumed by downstream code auditors. Compared with classical NLP-based approaches, the main advantage of recent LLM-based systems is their ability to handle the semantic diversity of real documentation, reducing reliance on task-specific semantic templates.

4.2.2 Protocol Specifications

Protocol specifications can be viewed as a special form of documentation. Compared with general software documentation, protocol specifications are domain-specific standards that are usually more normative, structured, and authoritative, and are most commonly published as RFCs (Request for Comments). Typically, they prescribe protocol behavior, including message formats, state transitions, error handling, and protocol-level constraints that conforming implementations are expected to satisfy. As a result, protocol specifications provide a particularly valuable semantic source for auditing protocol implementations. Once extracted into machine-readable forms, they can serve as input grammars, state machines, conformance oracles, or protocol-specific constraints for code auditing.

Despite their relatively standardized form, protocol specifications remain natural-language documents. Their requirements are often distributed across long documents through section hierarchies, cross-references, conditional statements, and normative keywords such as “MUST” and “SHOULD”. Understanding whether an implementation conforms to a protocol therefore requires more than sentence-level interpretation. An auditor must recover how message formats, field constraints, and stateful interaction rules compose into coherent protocol-level behavior. Traditionally, this process has relied heavily on manual reasoning by domain experts, making specification analysis time-consuming, error-prone, and difficult to scale.

Recent work shows that LLMs provide a promising new avenue for extracting machine-readable protocol knowledge from natural-language standards. Existing LLM-based approaches can be roughly divided according to the source from which protocol knowledge is obtained. Several approaches elicit latent protocol knowledge from pretrained models, while others explicitly ground the model’s reasoning in RFC documents and extract protocol constraints, message formats, or state machines from the standard text.

A representative example of the first line is ChatAFL, an LLM-guided fuzzer for network protocol implementations (Meng *et al.*, 2024b). ChatAFL is motivated by the observation that modern LLMs are

pretrained on large corpora that include substantial protocol-related text and therefore may implicitly encode knowledge about message formats and valid interaction patterns. For a given text-based protocol, such as RTSP or SMTP, ChatAFL queries LLMs to generate message grammars describing the structure of protocol messages and uses these grammars to guide mutation-based fuzzing. It also uses LLMs to enrich the initial seed corpus by increasing the diversity of recorded message sequences. However, ChatAFL does not synthesize a complete protocol state machine directly with LLMs. Instead, it retains the execution-feedback loop and uses LLMs only as local sources of guidance. When fuzzing enters a coverage plateau, ChatAFL prompts an LLM with the recent client-server communication history and asks it to generate the next client request that may induce a new state transition. The generated request is then inserted into the current message sequence and executed against the implementation. If the resulting sequence increases coverage, ChatAFL uses the server response and state feedback to update the inferred protocol state machine. This design illustrates how pretrained protocol knowledge can be useful for fuzzing, while also suggesting a limitation of such knowledge-only grounding: the model may omit or invent protocol details when the target protocol is underrepresented in pretraining data. More broadly, although ChatAFL is designed for protocol fuzzing, it suggests that protocol knowledge acquired during pretraining can provide useful semantic hints for downstream program analysis tasks, including bug detection, when combined with additional validation mechanisms.

The second line of work targets network protocols whose specifications are well documented in RFC documents. Specifically, these approaches ground extraction in the authoritative specification and aim to recover protocol constraints, message formats, state machines, or other structured artifacts from RFC documents. They differ in the kinds of artifacts they extract and in how these artifacts support auditing. Several systems focus on message formats and parser-level constraints. For example, 3DGen (Fakhoury *et al.*, 2025) uses specification knowledge to support the generation of binary-format parsers, while ParCleanse (Zheng *et al.*, 2025b) and ParVAL (Zheng *et al.*, 2025c) use RFC-derived information to validate network protocol parsers. Although

these techniques differ in their specific targets, they share a common goal of converting natural-language protocol standards into structured knowledge that can assist verified code generation, program verification, and bug detection.

Grounding LLMs in RFC documents improves fidelity to the standard, but it also introduces a new challenge. Concretely, protocol specifications are often long, internally cross-referenced, and globally constrained. Context-window limitations and context-collapse effects can cause LLMs to miss important constraints or apply them inconsistently across sections. Existing approaches therefore commonly adopt divide-and-conquer strategies that decompose specifications into smaller, structurally coherent units aligned with section hierarchies. For example, ParCleanse constructs a document tree to compose message formats extracted from individual sections (Zheng *et al.*, 2025b), whereas AutoSpec performs section-wise behavioral extraction and deterministically merges extracted states, transitions, and constraints into a protocol-wide state-transition graph (Liu *et al.*, 2025). These designs show that the central difficulty is not merely extracting isolated facts from RFCs, but maintaining cross-section and cross-state consistency when converting long standards into auditable protocol models.

Discussion. The analysis of semi-structured protocol documentation has a structural analogy to call graph analysis in classical static analysis. The analogy is not that document sections behave like program functions, but that both settings require recovering a graph-structured representation from distributed local evidence. In static analysis, call graph analysis links call sites and potential callees to support interprocedural reasoning. In protocol-specification analysis, references among terms, fields, sections, and state descriptions can similarly be used to construct a document-level graph. Such a graph helps organize long documents, propagate constraints across related sections, and support consistent extraction of protocol-level semantics.

4.2.3 Bug Knowledge Sources

Beyond developer-facing documentation and protocol standards, LLM-based code auditing can also exploit a broader class of *bug knowledge sources*. These sources are produced during bug classification, disclosure, discussion, and repair. Unlike library documentation or protocol specifications, they are usually not authoritative descriptions of how a system should behave. Instead, they encode how software bugs are conceptualized, observed, explained, and fixed in practice. As a result, they provide complementary semantic evidence for code auditing, which helps define what counts as a bug, suggests what program behaviors are security-relevant, or supplies concrete examples from which reusable bug patterns can be learned.

Bug knowledge sources differ in their level of curation and proximity to code. Curated bug taxonomies and disclosure records, such as CWE (CWE, 2025) and CVE (CVE, 2025), provide high-level security concepts and concrete bug instances. Bug reports, commits, and patches provide code-grounded evidence of how defects manifest and how developers repair them. Online technical discussions provide community knowledge about API usage, debugging practices, and security-relevant edge cases. These sources therefore support different forms of grounding for LLM-aided code auditors.

Curated bug records, such as the Common Weakness Enumeration (CWE) and Common Vulnerabilities and Exposures (CVE), are widely used sources of curated bug knowledge (CWE, 2025; CVE, 2025). CWE organizes weaknesses into abstract classes, often describing their causes, consequences, examples, and mitigations, whereas CVE assigns identifiers and descriptions to concrete publicly disclosed bugs. Together, they form a bridge between high-level bug concepts and real-world failure instances. One line of work uses CWE/CVE-style knowledge for proactive bug discovery. In this setting, the goal is not merely to classify code by a known weakness label, but to convert bug knowledge into specifications or analysis guidance that can be applied to unseen code. For example, IRIS (Li *et al.*, 2025f) takes a target project and a CWE class as input, uses LLMs to infer CWE-specific taint specifications such as sources and sinks, and integrates the inferred specifications with CodeQL for

codebase-wide analysis. Vul-RAG (Du *et al.*, 2026) constructs a knowledge base from existing bug instances by distilling information about functionality, bug causes, and fixes, and retrieves relevant knowledge to guide LLM-based bug detection. These systems illustrate a common design principle that curated bug records are not used as final answers, but are transformed into intermediate specifications, retrieved knowledge, or reasoning instructions that ground downstream code analysis. Another line of work focuses on post-disclosure bug localization. Here, the bug is already known, and the auditing task is to identify where the disclosed bug appears in the corresponding codebase. CVE descriptions are therefore treated as natural-language bug specifications that must be aligned with source code. VFFinder (Wu *et al.*, 2024), for example, uses CVE descriptions and source code to identify buggy functions. This setting emphasizes that the auditor must connect a bug report to the concrete functions or code regions that realize the root cause.

Project histories provide a localized form of security evidence for code auditing. They include issue reports, commit records, code review discussions, and bug-fixing patches that document how security-relevant defects have appeared, been diagnosed, and been repaired within a particular codebase or program domain. Unlike CWE/CVE records, which organize bug knowledge across projects and software ecosystems, project histories preserve local development experience. Their value therefore lies not only in individual bug-fix examples, but also in the recurring patterns exposed by a sequence of related development activities. In the same project or domain, bugs often recur because related code shares similar abstractions, implementation conventions, and operational assumptions. Project histories can therefore act as localized bug memory, enabling auditors to recognize recurring defect patterns that may not be fully captured by general bug taxonomies.

A common strategy is to learn bug patterns from representative historical examples and generalize them to new code within the same project or program domain. BugScope (Guo *et al.*, 2025a) follows this intuition by emulating how human auditors learn from buggy and non-buggy examples. Given representative examples, it synthesizes a retrieval strategy to collect relevant program context and constructs a

tailored detection prompt for LLM-based auditing. This design treats buggy and non-buggy examples not simply as isolated labels, but as demonstrations of domain-specific bug patterns that can guide the search for similar defects. A complementary strategy is to distill patch knowledge into reusable symbolic detectors. KNighter (Yang *et al.*, 2025a) synthesizes static analysis checkers from historical patches in the Linux kernel. Rather than asking an LLM to directly scan the entire Linux kernel codebase, KNighter uses LLMs to extract a bug pattern from patches, generate a specialized checker, and validate or refine the checker against the original patch context. This design converts patch-derived knowledge into a scalable static-analysis artifact, allowing it to scale to large codebases.

Community technical discussions, including Stack Overflow questions, project mailing lists, issue discussions, and technical forums, provide a less curated but highly practical source of software knowledge. These discussions often capture developer-facing experience that may be absent from official documentation, such as common API misuse patterns, debugging strategies, and workarounds for version-specific behavior. For code auditing, such knowledge can provide useful hypotheses about where defects arise and what contextual assumptions developers commonly miss. However, community discussions are weaker evidence than curated bug records or code patches. They may be outdated, incomplete, contradictory, or specific to a particular environment. LLMs pretrained on web-scale corpora may implicitly encode parts of this community knowledge, but because it is stored latently in model parameters, it is difficult to attribute, update, or verify. Recent systems therefore increasingly favor retrieval-augmented designs that ground generated answers in explicitly retrieved community-authored content. For example, Stack Overflow’s AI Assist (Sajor, 2025) uses a Retrieval-Augmented Generation (RAG) pipeline over Stack Overflow and Stack Exchange content, prioritizing human-validated answers and providing attribution. Although such systems are not themselves code auditors, they illustrate a broader grounding pattern that is relevant to auditing, highlighting that community knowledge is most useful when retrieved, attributed, and treated as supporting evidence rather than as unverified

model memory.

Discussion. Bug knowledge sources differ from documentation-based sources in both authority and evidential structure. Library documentation and protocol specifications describe intended behavior, whereas bug records, patches, and discussions describe security knowledge accumulated through failures, repairs, and community experience. This makes them valuable but also noisy. The central design question is therefore how to turn heterogeneous bug evidence into auditable artifacts, such as bug specifications, detection prompts, static-analysis checkers, or localization queries. Across the systems above, LLMs are most reliable when their outputs are grounded in curated records, validated against code or patches, linked to retrieved evidence, or compiled into analyzable artifacts. In this sense, the role of LLMs is not to replace security expertise, but to operationalize distributed bug knowledge so that downstream auditors can use it systematically.

4.3 Summary

This chapter surveys LLM-driven primitive analyses over both code and non-code artifacts. Within the agentic view developed in Chapter 3, these analyses correspond to the primitive tools that transform retrieved artifacts into intermediate evidence for auditing. Their outputs are not final bug reports, but formal/informal semantic specifications. In this sense, primitive analysis provides the bridge between raw software artifacts and the higher-level auditing decisions studied in the next chapter.

For code artifacts, the surveyed techniques show that LLMs are most useful at semantic bottlenecks where classical static analyses are costly, brittle, or incomplete. Pointer analysis and dependency analysis recover memory-related and data-flow facts. Call graph analysis exposes interprocedural structure needed for cross-function reasoning. Invariant generation summarizes stable properties of program states, while execution trace prediction reveals path-specific runtime behavior. These primitive facts directly support the bug classes discussed in Chapter 2. For example, memory bugs require aliasing, lifetime, and nullability

facts. Taint-style bugs require data-flow, call-graph, and specification facts. Numeric bugs often rely on value invariants and path conditions. Functional bugs depend on recovering intended behavior. However, LLM-derived code facts remain fallible. They can reduce dependence on full build environments and often operate without structured IRs, but they do not provide the soundness, completeness, or consistency ensured by classical static analysis. As a result, the LLM-driven primitive code analyses demonstrate that LLMs are most effective when used to approximate, prioritize, or refine semantic properties, rather than to replace classical static analyses wholesale.

For non-code artifacts, LLMs address the complementary side of the semantic barrier problem. Many auditing-relevant specifications are not present in code at all, but are distributed across library documentation, protocol standards, bug records, project histories, patches, and community discussions. The systems surveyed in this chapter use LLMs to convert such informal or semi-structured evidence into auditable artifacts, which can be viewed as a form of externalized long-term auditing knowledge, i.e., once extracted, they help auditors interpret code behavior, instantiate bug-specific analyses, and reuse knowledge across different projects, versions, and auditing tasks.

Taken together, the primitive analyses surveyed in this chapter are composed and orchestrated according to the demands of specific auditing subtasks. Rather than committing to static analysis tools or LLMs only, practical auditors dynamically combine static tools and LLM-driven reasoning, switching between static analysis tools and semantic intuition offered by LLMs. The next chapter will build on this foundation by examining how these primitive analyses are instantiated and coordinated in neurosymbolic auditing systems, illustrating how they collectively enable scalable code auditing across real-world codebases.

5

Auditing Code as Human Auditors

As discussed in Chapter 2, classical static analysis faces fundamental challenges in code auditing due to IR reliance, limited customizability, and the semantic barrier. To address these challenges, Chapter 3 introduces the agent-centric perspective, highlighting how human auditors use various tools, software artifacts in the auditing environment, and memory to reason effectively about code behavior. Chapter 4 further elaborates on LLM-driven primitive analyses, demonstrating their capability to analyze both code and non-code artifacts.

Building upon these foundations, this chapter explores how neurosymbolic approaches leverage LLMs and symbolic analysis tools, including classical static analyzers, to closely simulate human auditing behavior. As shown in Figure 5.1, we start from two typical modes of human code auditing, namely manual code auditing and semi-automatic code auditing, and examine how existing approaches effectively simulate these two auditing modes in Section 5.1 and Section 5.2, respectively. Beyond these academic efforts, we also systematically discuss recent industrial advances in code auditing in Section 5.3 before summarizing this chapter in Section 5.4.

