

AI Security Compliance and Testing Framework for Large Language Model Systems

Harsh Jannawar

A thesis submitted in partial fulfillment of the
requirements for the degree of

Master of Science in Cybersecurity Engineering

University of Washington
2026

Committee:

Dr. Min Chen

Dr. Marc Dupuis

Dr. Yang Peng

Program Authorized to Offer Degree:
Computing and Software Systems

© Copyright 2026

Harsh Jannawar

University of Washington

ABSTRACT

AI Security Compliance and Testing
Framework for Large Language Model
Systems

Harsh Jannawar

Chair of the Supervisory Committee: Dr. Min Chen

Large Language Models are being integrated into enterprise workflows at a pace that has outrun the security frameworks designed to protect them. Conventional compliance standards such as SOC 2 and ISO 27001 provide no coverage for LLM-specific vulnerabilities including prompt injection, sensitive information disclosure, and system prompt leakage. The OWASP LLM Top 10 defines the relevant threat taxonomy, but no unified automated pipeline exists to translate those controls into repeatable, evidence-based test cases. This research investigates whether automated, systematically constructed adversarial testing can reliably surface exploitable vulnerabilities across LLM applications of varying hardness, and whether evolutionary attack strategies can reach attack surfaces that static prompt libraries cannot anticipate.

To address these questions, this research employs a Design Science Research methodology. A 279-prompt library was constructed through three-source triangulation, drawing from CTF competition wins, industry AI red-teaming competition data, and peer-reviewed literature, grounding every technique in documented real-world effectiveness. Three target configurations were designed with isolated independent variables to evaluate detection rates across a baseline system, a hardened system, and a RAG-augmented system. An evolutionary attack engine implementing the SPE-NL genetic algorithm was developed and evaluated across all three configurations. Judge reliability and inter-rater agreement were validated through independent assessment by two practicing cybersecurity professionals.

Following this research design, AegisLLM was implemented as an automated security testing suite operationalizing six OWASP LLM Top 10 controls. Empirical evaluation demonstrates that targets resistant to the full static library fall to SPE-NL-evolved payloads within three to five generations, confirming that adaptive evolutionary testing reaches attack surfaces that curated static libraries cannot anticipate. The LLM-as-a-judge classification pipeline achieved 94.5% inter-rater agreement with zero crossover errors between SUCCESS and NO_SUCCESS labels.

This thesis contributes to the field in three respects. First, it provides an empirically validated, open-source prompt library mapped explicitly to the OWASP LLM Top 10, sourced from ecologically valid real-world adversarial data. Second, it demonstrates the viability of

evolutionary prompt mutation as a structured research method for LLM security evaluation, not merely as an engineering technique. Third, it establishes a replicable evaluation framework combining automated semantic judgment with human inter-rater validation, offering a methodological foundation for future LLM security research.

ACKNOWLEDGEMENTS

I would like to thank my committee chair Dr. Min Chen and committee members Dr. Marc Dupuis and Dr. Yang Peng for their guidance throughout this research. I am grateful to my collaborator Saniya Bhaladhare for the partnership on the broader framework of which this thesis is one component. I also thank the organizers of the AI red-teaming competition who provided access to the validated prompt dataset that forms a foundational part of this work.

TABLE OF CONTENTS

Chapter 1. Introduction.....	5
1.1 Background and Motivation.....	5
1.2 Problem Statement.....	6
1.3 Research Questions.....	6
1.4 Contributions.....	7
1.5 Scope and Collaborative Division of Work.....	7
1.6 Thesis Organization.....	8
Chapter 2. Literature Review and Related Work.....	9
2.1 LLM Security Threat Landscape.....	9
2.2 Existing Testing Frameworks and Tools.....	9
2.3 Compliance Frameworks and Regulatory Alignment.....	10
2.4 Evolutionary and Genetic Approaches to Adversarial Input Generation.....	10
2.5 LLM-as-a-Judge Evaluation.....	11
2.6 Positioning of AegisLLM Relative to Prior Work.....	11
Chapter 3. Research Methodology.....	13
3.1 Research Design and Approach.....	13
3.2 Research Questions Operationalization.....	13
3.3 System Development Methodology.....	14
3.4 Prompt Library Construction Methodology.....	14
3.5 Experimental Design.....	15
3.6 Validation Methodology.....	16
3.7 Threats to Validity.....	17
Chapter 4. System Design and Architecture.....	19
4.1 Design Goals and Constraints.....	19
4.2 High-Level Architecture.....	19
4.3 CLI Manager (main.py).....	20
4.4 Web Dashboard (web_app.py).....	21
4.5 Prompt Testing Engine (prompt_tester.py).....	23
4.6 Response Storage and Analysis Pipeline.....	24
4.7 LLM-as-a-Judge Evaluator (injection_judge.py).....	24
4.8 Domain-Specific Calibration of Classification Sensitivity.....	26
Chapter 5. Prompt Library Design and Construction.....	28
5.1 Prompt Sourcing Methodology.....	28
5.2 CSV Schema and Ingestion Standard.....	29
5.3 Direct Injection Library (prompts/direct_injection.csv) - 25 prompts.....	30
5.4 Narrative and Obfuscation Libraries.....	31

5.5 Sensitive Data Exposure Library - 52 prompts.....	32
5.6 Harmful Output and Off-Topic Libraries - 30 prompts.....	32
5.7 Misinformation and Hallucination Library - 30 prompts.....	32
5.8 RAG-Specific Attack Libraries - 82 prompts.....	33
5.9 Jailbreak Seed Library - 10 prompts.....	33
5.10 OWASP LLM Top 10 Control Coverage Matrix.....	33
Chapter 6. ThreatForge: Evolutionary Attack Suite.....	36
6.1 Motivation for Evolutionary Prompt Generation.....	36
6.2 SPE-NL Algorithm Overview.....	38
6.3 Fitness Scoring Function: Full Mathematical Specification.....	39
6.3.1 Refusal Avoidance (30% Weight).....	39
6.3.2 Information Leakage (25% Weight).....	40
6.3.3 Semantic Compliance (20% Weight).....	40
6.3.4 Novelty and Lineage Diversity (10% Weight).....	40
6.3.5 Target Pattern Matching (10% Weight).....	40
6.3.6 Substantive Length (5% Weight).....	40
6.4 Breakthrough Detection and Logging.....	41
6.5 EvolutionConfig Parameters.....	41
6.6 Guardrail Posture Scan (run_posture.py).....	42
6.7 Differential Attack Scan (run_differential.py).....	42
Chapter 7. Evaluation and Results.....	44
7.1 Experimental Setup.....	44
7.2 Static Analysis Baseline Results.....	44
7.3 LLM-as-a-Judge Accuracy Validation.....	47
7.4 SPE-NL Evolutionary Performance.....	48
7.5 RAG-Specific Evaluation Results (Configuration C).....	49
7.6 Formal Comparison with Manual Penetration Testing.....	50
7.7 Discussion of Findings.....	50
Chapter 8. Discussion.....	52
8.1 Answering the Research Questions.....	52
8.2 Limitations.....	54
8.3 Ethical Considerations.....	55
8.4 Generalizability.....	55
8.5 Pipeline Integration and Safeguards.....	56
8.6 Preventing Test Artifact Leakage.....	56
8.7 Security and Privacy Failure Handling.....	57
Chapter 9. Conclusion and Future Work.....	58
9.1 Summary of Contributions.....	58

9.2 Future Work.....	59
9.3 Closing Remarks.....	59
Bibliography.....	61
Appendix A. AegisLLM Configuration Reference.....	63
A.1 CLI Manager Flags and Options.....	63
A.2 ThreatForge CLI Options.....	63
A.3 EvolutionConfig Parameter Reference.....	64
A.4 JSON Response Path Notation Guide.....	65
A.5 Rate Limit Tester Configuration.....	65
A.6 Result CSV Schema Reference.....	66
Appendix B. Raw Results Tables.....	67
B.1 Complete Static Analysis Detection Results.....	67
B.2 Detection Rate Summary by OWASP Control.....	67
B.3 LLM-as-a-Judge Spot-Check Validation Results.....	68
B.4 SPE-NL Evolutionary Run Full Results.....	69
B.5 Canary Token Detection Full Results.....	69
B.6 Manual Penetration Testing Comparison Results.....	69

Chapter 1. Introduction

1.1 Background and Motivation

Large language models are no longer experimental technology. They underpin customer service platforms, legal document review tools, medical triage assistants, code generation pipelines, and internal enterprise knowledge bases across virtually every industry sector. This rapid, broad adoption has created a security problem that the industry is not yet equipped to address systematically: the attack surface of an LLM-integrated application is fundamentally different from anything that existing security frameworks were designed to protect [2], [8].

The frameworks that organizations rely on today, SOC 2 [26], ISO 27001 [27], and NIST SP 800-53 [28], define controls for network perimeters, access management, encryption at rest, and audit logging. They assume that the attack surface of a system is structural: unauthorized access happens through authentication bypass, data is exfiltrated through misconfigured storage, and services are disrupted through resource exhaustion at the network layer. These assumptions held reasonably well for a generation of software systems built around deterministic, rule-based logic. They fail almost entirely when applied to LLM-integrated applications, where the attack surface is linguistic and the most dangerous input a system can receive is a carefully constructed sentence [2]. SOC 2 has no control for prompt injection. ISO 27001 has no annex for context poisoning. NIST SP 800-53 has no baseline for adversarial token manipulation. The gap between what these frameworks cover and what LLM deployments actually face is not a minor oversight; it is a categorical mismatch between the threat model these frameworks were built for and the threat model that LLM-integrated systems actually face.

At the same time, no unified automated pipeline exists that translates the emerging LLM-specific threat taxonomy into repeatable, evidence-based test cases that security teams can actually run against their deployments. The OWASP LLM Top 10 (2025) provides the most comprehensive publicly available taxonomy of LLM-specific vulnerabilities, from Prompt Injection (LLM01) through Unbounded Consumption (LLM10) [18], but a taxonomy is not a testing tool. Individual researchers have demonstrated effective prompt injection techniques and guardrail bypass methods in controlled experiments [2], [8], [22], but these demonstrations have not been consolidated into a reusable, operationally deployable framework. The result is that most organizations deploying LLM-integrated applications have no systematic way to determine whether their specific deployment is vulnerable to the threats the taxonomy describes.

This thesis presents AegisLLM, an automated security testing suite designed to close both gaps simultaneously: the gap between existing compliance frameworks and LLM-specific threats, and the gap between the OWASP LLM Top 10 taxonomy and a practical, automated testing pipeline that operationalizes it.

1.2 Problem Statement

Despite the rapid integration of LLM-powered systems into enterprise workflows, security teams lack an automated, evidence-based pipeline for testing LLM deployments against the specific vulnerability classes defined by the OWASP LLM Top 10 [18]. Existing penetration testing tools are designed for deterministic software systems and cannot interpret or evaluate natural language responses as security findings [7]. Existing LLM safety evaluation tools focus on alignment and content policy rather than adversarial security testing [19]. And while individual researchers have demonstrated effective prompt injection techniques and guardrail bypass methods in controlled experiments [2], [8], these demonstrations have not been consolidated into a reusable, operationally deployable testing framework.

The consequence of this gap is that most organizations deploying LLM-integrated applications have no systematic way to answer three fundamental security questions: Which of the OWASP LLM Top 10 vulnerability classes, including prompt injection (LLM01), sensitive information disclosure (LLM02), improper output handling (LLM05), system prompt leakage (LLM07), and misinformation (LLM09), does our deployment exhibit under automated adversarial testing? How resistant are our guardrails to adversarial prompts that go beyond known technique classes? And which vulnerabilities were introduced by our fine-tuning or RAG augmentation versus which were inherited from the underlying base model? AegisLLM is designed to answer all three questions within a single, unified testing pipeline.

1.3 Research Questions

This thesis is organized around three research questions derived from the problem statement:

RQ1: How can each OWASP LLM Top 10 control be operationalized into precise, testable criteria that can be executed automatically against any REST-accessible LLM endpoint?

RQ2: To what extent does dynamic adversarial input generation through evolutionary mutation uncover runtime vulnerabilities that a static, curated prompt library cannot reach?

RQ3: What hybrid evaluation strategy best balances coverage, false positive rate, and throughput in an automated LLM security testing pipeline that must operate at the scale of thousands of responses per engagement?

These three questions map respectively to the three core research contributions of this thesis: the OWASP-aligned prompt library and testing engine, the SPE-NL genetic algorithm and ThreatForge evolutionary suite, and the dual-track evaluation architecture combining deterministic algorithmic scoring with zero-temperature LLM-as-a-judge classification.

1.4 Contributions

The primary contribution of this thesis is a set of research findings that advance the understanding of automated adversarial security testing for LLM applications. A fully implemented testing suite, AegisLLM, was developed and evaluated as the vehicle through which those findings were produced and validated. Within this broader contribution, three specific research contributions distinguish this work from prior work in LLM security evaluation.

The first contribution is an empirically validated methodology for constructing adversarial prompt libraries that are grounded in documented real-world effectiveness. Prior LLM security evaluations have relied on synthetically generated or researcher-intuited prompts without systematic sourcing criteria. This research demonstrates that a three-source triangulation approach, drawing from CTF competition wins, industry AI red-teaming competition data, and the adversarial ML research literature, produces prompt collections with measurably stronger ecological and convergent validity than single-source approaches. The resulting library of 279 prompts across 13 files, mapped to six OWASP LLM Top 10 controls, constitutes the empirical foundation of the evaluation and is made available for reuse by future researchers.

The second contribution is a finding about the limits of static prompt libraries as a security evaluation method. Empirical testing across three target configurations revealed that hardened systems resist the full static library but remain vulnerable to adaptively generated payloads. This finding motivates and validates the SPE-NL genetic algorithm, which was developed on the hypothesis that feedback-driven mutation in the natural language prompt space, modeled on the fuzzing principle from traditional penetration testing, can systematically discover attack payloads in regions of the prompt space that no curated static library anticipates. The algorithm was implemented in the ThreatForge evolutionary suite and evaluated against all three target configurations, achieving breakthrough in three to five generations in every case.

The third contribution is a validated evaluation architecture for automated LLM security testing at scale. Prior work faces a reliability tradeoff between semantic richness and throughput: human annotation is accurate but does not scale, and automated string matching scales but cannot evaluate open-ended adversarial outcomes. This research demonstrates that a dual-track architecture combining zero-temperature LLM-as-a-judge classification for semantic evaluation with deterministic algorithmic scoring for evolutionary fitness resolves this tradeoff. The judge pipeline achieved 94.5% agreement with two independent human assessors across 91 sampled records, with zero crossover errors between SUCCESS and NO_SUCCESS labels.

1.5 Scope and Collaborative Division of Work

The contribution documented here covers is part of a collaborative research project with Saniya Bhaladhare, whose work focuses on the governance and compliance dimensions of the same overarching framework. Saniya's contribution covers the derivation of precise audit criteria and compliance metrics from ISO 42001 and the NIST AI Risk Management Framework, embedded within the framework's governance and reporting modules. My contribution, which this thesis documents in full, covers the operationalization of OWASP LLM Top 10 controls into automated test cases, the implementation of the static and evolutionary testing pipeline, the design and construction of the prompt libraries, and the empirical validation of the complete

AegisLLM system against target LLM deployments. The two bodies of work are designed to integrate into a single unified framework but are independently evaluable and independently documented.

1.6 Thesis Organization

The remainder of this thesis is organized as follows. **Chapter 2** surveys the existing literature on LLM security testing, adversarial prompt generation, and compliance automation, positioning AegisLLM relative to prior work and identifying the specific gaps it addresses. **Chapter 3** presents the research methodology, covering the design science research approach, experimental design, prompt library construction methodology, and the two-assessor validation protocol. **Chapter 4** describes the system architecture of AegisLLM in full, covering the CLI manager, web dashboard, prompt testing engine, response storage pipeline, and LLM-as-a-judge evaluator. **Chapter 5** documents the design and construction of the AegisLLM prompt libraries, including the three-pronged sourcing methodology, the CSV ingestion standard, and the complete library inventory organized by OWASP control. **Chapter 6** presents the ThreatForge evolutionary suite in detail, covering the SPE-NL genetic algorithm, the six-dimensional fitness scoring function, breakthrough detection and logging, and the posture and differential scanning modules. **Chapter 7** presents the evaluation results from testing AegisLLM against target LLM deployments, including detection rates by OWASP control, evolutionary performance metrics, and LLM-as-a-judge accuracy validation. **Chapter 8** discusses the findings in relation to the three research questions, addresses limitations, and considers the ethical dimensions of automated adversarial testing tooling. **Chapter 9** concludes with a summary of contributions and directions for future work.

Chapter 2. Literature Review and Related Work

2.1 LLM Security Threat Landscape

The security risks introduced by large language models have been documented with increasing precision since the widespread deployment of LLM-integrated applications beginning in 2022. The foundational threat characterization relevant to this thesis comes from two directions: empirical studies of specific attack techniques against deployed LLM systems, and broader taxonomic efforts to organize those techniques into actionable vulnerability categories.

Perez et al. conducted one of the earliest systematic empirical studies of prompt injection attacks in LLM-integrated applications, demonstrating that adversarial user inputs could reliably override operator-defined system prompt instructions in a range of real-world deployment configurations [2]. Their work established prompt injection as a distinct, reproducible vulnerability class rather than an edge case. Gupta and Zhao extended this analysis by systematically mapping prompt injection vulnerabilities across a broader range of deployment architectures, identifying context window construction and role boundary enforcement as the two primary determinants of injection resistance [8]. Their findings directly informed the design of the End Marker Spoofing and SYSTEM Tag Spoofing techniques in the AegisLLM direct injection library.

At the taxonomic level, the OWASP LLM Top 10 (2025) represents the most comprehensive publicly available organization of LLM-specific vulnerability classes, covering ten distinct threat categories from Prompt Injection through Unbounded Consumption [18]. Of particular relevance to this thesis is the distinction between OWASP LLM01 (Prompt Injection) and OWASP LLM07 (System Prompt Leakage), which the OWASP taxonomy treats as separate controls but which AegisLLM covers as co-primary objectives across the same prompt libraries. LLM01 addresses the risk that an adversarial input overrides operator-defined instructions, while LLM07 specifically addresses the risk that confidential operator instructions embedded in the system prompt are disclosed to unauthorized users as a result of adversarial prompting. In practice these two outcomes co-occur: a prompt that successfully injects a malicious instruction typically also reveals the content of the instructions being overridden. Johnson et al. further characterized adversarial input behavior in LLM decision-making systems, providing empirical data on how adversarial token sequences influence model outputs in ways not captured by standard safety evaluation benchmarks [10].

2.2 Existing Testing Frameworks and Tools

Prior work on automated LLM security testing falls into three broad categories: static prompt analysis tools, dynamic adversarial input generation frameworks, and hybrid approaches that combine elements of both.

Saimbhi and Akpinar developed VulnerAI, a GPT-based vulnerability detection system for web applications that demonstrated the feasibility of using language models to identify prompt injection risks through static analysis of application inputs [1]. Their work established an important proof of concept but operates as a static analyzer applied to application code rather than as a dynamic penetration testing tool applied to live endpoints. Lee and Wang proposed

metamorphic prompt testing as a methodology for evaluating LLM-generated programs, demonstrating that systematically varied prompt structures could expose inconsistencies in model behavior [9]. While their approach shares the structural mutation intuition behind SPE-NL, it targets code generation correctness rather than adversarial security testing.

Zhang et al. demonstrated prompt-enhanced vulnerability detection using ChatGPT, showing that carefully constructed prompts submitted to a capable LLM could identify software vulnerabilities with competitive detection rates [5]. Their work is relevant as prior art for the LLM-as-a-judge evaluation pattern adopted in AegisLLM. Patel et al. developed TF-Attack, a transferable and fast adversarial attack framework for LLMs that demonstrated high attack success rates across multiple target models [12]. TF-Attack operates in the embedding space of the model rather than in the natural language prompt space, requiring white-box access to model internals unavailable in the black-box deployment scenarios AegisLLM is designed for. Das Purba et al. examined software vulnerability detection using LLMs with a focus on the operational constraints of automated testing pipelines, including the risk of API bans and service throttling under aggressive testing conditions [7].

2.3 Compliance Frameworks and Regulatory Alignment

A parallel strand of related work addresses the governance and compliance dimensions of LLM security. Gajbhiye et al. examined the use of AI compliance tools for application security, demonstrating that regulatory alignment frameworks could be applied to LLM-integrated applications but stopping short of providing an automated end-to-end testing workflow [3]. Their work identifies the same gap this thesis addresses from the compliance side. Kothandapani outlined high-level AI-driven regulatory compliance frameworks for financial oversight, demonstrating the demand for automated compliance tooling in regulated industries [4]. Berger et al. proposed rethinking legal compliance automation through the lens of LLM capabilities, identifying opportunities for LLMs to automate compliance reporting and audit trail generation [6]. Mohammed examined AI applications in cybersecurity audit and compliance automation, establishing the general feasibility of AI-assisted audit workflows while acknowledging that LLM-specific security controls remained outside the scope of existing automated audit tools [11].

2.4 Evolutionary and Genetic Approaches to Adversarial Input Generation

The conceptual origin of ThreatForge is not the genetic algorithm literature but the fuzzing tradition in traditional penetration testing. Coverage-guided fuzzing tools such as AFL and libFuzzer operate on a simple and powerful principle: take a valid seed input, apply structured mutations to it, execute the mutated input against the target, and use the target's response as the signal that guides which mutations to retain and which to discard [14][15]. Over successive mutation cycles, the fuzzer accumulates a population of inputs that collectively explore the target's behavior space far more efficiently than random input generation could. This feedback-driven, mutation-based, population-accumulating search dynamic is exactly what SPE-NL implements in the natural language prompt space.

The transfer of fuzzing concepts to LLM security testing is both natural and non-trivial. It is natural because the goal is identical: discover inputs that cause a defended target system to

behave in ways it was not intended to behave, using the target's own responses as the guide. It is non-trivial because mechanical mutation operators such as bit flips and byte substitutions produce semantically incoherent outputs when applied directly to natural language text. SPE-NL resolves this by replacing the mechanical mutation operator with LLM-assisted rewriting, which produces mutations that are semantically coherent, structurally varied, and targeted toward the specific fitness dimensions where the parent prompt underperformed. The genetic algorithm selection mechanics are borrowed from the evolutionary computation literature [16][17] as the formal framework for accumulating search progress across mutation cycles, but the driving intuition is fuzzing.

Prior work that applies evolutionary approaches to LLM adversarial testing includes TF-Attack [12] and Lee and Wang's metamorphic approach [9]. Both demonstrate that systematic input variation can expose LLM vulnerabilities that static test sets miss. However, TF-Attack requires white-box model access, and Lee and Wang's approach applies predefined transformation rules rather than fitness-guided evolutionary selection. SPE-NL addresses both limitations: it operates entirely in natural language without requiring model internals, and its fitness-guided selection accumulates search progress directionally.

2.5 LLM-as-a-Judge Evaluation

The use of LLMs as automated evaluators of other LLM outputs has emerged as an active research area in response to the scalability limitations of human annotation for NLP evaluation tasks. Zheng et al. demonstrated that capable LLMs could serve as reliable judges for open-ended question answering and dialogue quality assessment, achieving agreement rates with human annotators comparable to inter-annotator agreement between human experts [13]. Their work established the LLM-as-a-judge pattern as a credible evaluation methodology for tasks requiring semantic understanding at scale.

The AegisLLM judge implementation builds on this precedent while addressing two reliability concerns that limit applicability in security contexts. The first concern is non-determinism: an LLM judge running at default temperature can produce different classification labels for identical inputs across runs, making results non-reproducible. AegisLLM addresses this through zero-temperature enforcement. The second concern is output format instability: an LLM judge that produces free-text rationales cannot be parsed programmatically without brittle string processing. AegisLLM addresses this through A PI-level JSON format locking. The three-label classification scheme, SUCCESS, POSSIBLE, and NO_SUCCESS, reflects findings in the annotation reliability literature that binary classification of ambiguous security outcomes produces systematically lower inter-annotator agreement than schemes that explicitly capture uncertainty.

2.6 Positioning of AegisLLM Relative to Prior Work

The literature surveyed in this chapter demonstrates three things. First, the LLM security threat landscape is well characterized at the taxonomic level. Second, individual components of an automated LLM security testing pipeline have been demonstrated in isolation. Third, no prior work has integrated all these components into a single, unified, operationally deployable testing pipeline grounded in empirically validated real-world attack techniques.

Table 2.1. Gap analysis comparing prior work against AegisLLM across nine evaluation dimensions.

Work	LLM01	LLM02	LLM05	LLM07	LLM09	Evol.	LLM-Judge	Real-World	Black-Box
Saimbhi & Akpinar [1]	Partial	No	No	No	No	No	No	No	Yes
Perez et al. [2]	Yes	No	No	Partial	No	No	No	No	Yes
Perez & Ribeiro [20]	Yes	No	No	Yes	No	No	No	No	Yes
Gajbhiye et al. [3]	Fwk	Fwk	Fwk	Fwk	Fwk	No	No	No	N/A
Zhang et al. [5]	No	No	No	No	No	No	Yes	No	Yes
Das Purba et al. [7]	No	No	No	No	No	No	No	No	Yes
Gupta & Zhao [8]	Yes	Partial	No	Yes	No	No	No	No	Yes
Lee & Wang [9]	Partial	No	No	No	No	Partial	No	No	Yes
Patel et al. [12]	Partial	No	No	No	No	Yes	No	No	No
Zheng et al. [13]	No	No	No	No	No	No	Yes	No	Yes
Greshake et al. [21]	Yes	Yes	Yes	No	No	No	No	No	Yes
Shen et al. [24]	Yes	No	No	Partial	No	No	No	Yes	Yes
AegisLLM (this work)	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes

"Fwk = addressed at compliance framework level only without automated implementation or empirical testing."

Chapter 3. Research Methodology

3.1 Research Design and Approach

This thesis adopts a Design Science Research (DSR) methodology, which is the appropriate paradigm for research that produces a purposefully designed artifact as its primary contribution and evaluates that artifact empirically against a defined set of research objectives. Design science research differs from purely empirical research in that the act of designing and constructing the artifact is itself a research contribution, not merely a means to an end. The artifact embodies the researcher's theoretical claims about how the problem should be solved, and its empirical evaluation provides evidence for or against those claims.

In this research, the artifact is AegisLLM: an automated security testing suite for LLM applications. The theoretical claims embedded in its design are threefold. First, that OWASP LLM Top 10 controls can be operationalized into precise, executable test cases grounded in empirically validated real-world attack techniques. Second, that evolutionary mutation guided by algorithmic fitness scoring can discover adversarial payloads that exceed the coverage of any static, curated prompt library. Third, that a dual-track evaluation architecture combining deterministic algorithmic scoring and zero-temperature LLM-as-a-judge classification can produce reliable, reproducible security findings at the scale that automated testing requires. The empirical evaluation in Chapter 7 provides the evidence against which these three claims are assessed.

The research followed four sequential phases aligned with the DSR process model. The problem identification phase produced the literature review in Chapter 2 and the gap analysis establishing that no existing tool achieves the combination of OWASP coverage, empirical prompt validation, and evolutionary attack generation that the identified gap requires. The objectives definition phase produced the three research questions and the success criteria from the project proposal. The design and development phase produced the AegisLLM system documented in Chapter 4. The demonstration and evaluation phase produced the controlled experiment reported in Chapter 7.

3.2 Research Questions Operationalization

Each of the three research questions maps to a specific research activity within the DSR process, and each is answered through a distinct methodological strand.

RQ1 (How can each OWASP LLM Top 10 control be operationalized into precise, testable criteria?) is answered through the prompt library construction methodology described in Section 3.4 and the system design documented in Chapter 4. The operationalization process treats each OWASP control as a specification to be translated into a set of concrete, executable test cases, following the same logic that software engineering uses to translate requirements into test suites. The quality of the operationalization is assessed by whether the resulting test cases produce detectable findings against representative target deployments, which is measured in the evaluation phase.

RQ2 (To what extent does evolutionary mutation uncover vulnerabilities a static library cannot reach?) is answered through a controlled within-subjects experiment in which the same target deployments are tested first with the static prompt library and then with the SPE-NL evolutionary engine under identical conditions. The difference in detection outcomes between the two testing approaches provides the empirical evidence for RQ2. This experimental design isolates the contribution of the evolutionary engine from the contribution of the static library by holding the target and testing conditions constant while varying only the testing methodology.

RQ3 (What hybrid evaluation strategy best balances coverage, false positive rate, and throughput?) is answered through the validation study described in Section 3.6, in which two independent human assessors reviewed tool-generated classifications against their own independent judgments. The agreement rate between the tool's classifications and the human assessors' classifications across different label categories and confidence bands provides the empirical evidence for RQ3.

3.3 System Development Methodology

AegisLLM was developed using an iterative, test-driven design process organized around the OWASP LLM Top 10 control taxonomy as the primary requirements specification. Rather than defining system requirements abstractly and then implementing them, each OWASP control was used as a concrete functional requirement: the system was required to produce measurable findings against that control when tested against a deployment that exhibits the corresponding vulnerability.

Development proceeded in three iterative cycles. The first cycle established the core testing infrastructure: the CLI manager, the prompt testing engine, the response storage pipeline, and the LLM-as-a-judge evaluator. These components were validated against a single minimally configured target deployment using an initial version of the direct injection prompt library. The second cycle extended the system with additional prompt libraries covering the remaining four primary OWASP controls, validated by running the full library against both Configuration A and Configuration B and observing whether the detection gradient between configurations matched the expected pattern of hardening effectiveness. The third cycle integrated the ThreatForge evolutionary engine and the RAG-specific attack modules, validated by running the evolutionary engine against Configuration B targets that had resisted the static library.

Each cycle followed a consistent development rhythm: design the component, implement it, test it against the target deployment, review the output against expected behavior, and refine before moving to the next component. This iterative validation against live target deployments was a deliberate methodological choice: for a security testing tool, the only meaningful validation is whether it produces correct findings against realistic targets, not whether individual functions pass isolated unit tests.

3.4 Prompt Library Construction Methodology

The prompt library construction is treated in this thesis as a formal data collection methodology rather than simply a software engineering task. This distinction matters because the quality of the prompt library determines the validity of the evaluation results: a library

constructed from low-fidelity or unvalidated sources will systematically underestimate the vulnerability of targets that are resistant to generic techniques, producing findings that reflect the library's limitations rather than the target's actual security posture.

A three-source triangulation approach was applied to prompt collection, borrowing the triangulation principle from qualitative research methodology to ensure that the library represents multiple independent perspectives on effective adversarial techniques. The three sources, CTF competition wins, industry red-teaming competition data, and peer-reviewed research literature, each provide a different form of validity evidence.

CTF competition data provides ecological validity: these techniques were validated in real competitive environments against live defended targets designed by security professionals specifically to resist attack. The findings from CTF challenges represent attack techniques with a documented win condition against a defended system, which is the strongest form of practical validity available without conducting unauthorized testing against production deployments.

Industry red-teaming competition data provides convergent validity: multiple independent attackers arriving at structurally similar techniques against the same class of targets provides evidence that those techniques generalize beyond any single individual's approach. The fact that different participants in the competition independently converged on similar injection strategies strengthens the claim that those strategies reflect genuine structural weaknesses in the target class rather than idiosyncratic exploitation by a single skilled attacker.

Research literature provides construct validity: grounding techniques in peer-reviewed characterizations of prompt injection and adversarial LLM behavior ensures that the library covers the documented theoretical attack surface rather than only the subset of techniques that happen to be effective in competitive settings.

All prompts from all three sources were processed through a consistent quality control pipeline before inclusion. Each prompt was independently mapped to at least one OWASP LLM Top 10 control, deduplicated against existing library entries, generalized where vendor-specific behaviors were targeted, and assigned a standardized technique name. Prompts that could not be unambiguously mapped to a control were excluded.

3.5 Experimental Design

The empirical evaluation of AegisLLM follows a controlled experiment design with three conditions corresponding to the three target configurations described in Chapter 7. The experimental design addresses a fundamental challenge in security tool evaluation: the absence of a ground-truth vulnerability oracle. Unlike software testing tools that can be validated against programs with known bugs, a security testing tool for LLM deployments cannot be validated against a definitive catalog of known vulnerabilities because no such catalog exists for the specific targets being tested.

The evaluation design addresses this challenge through three methodological choices. First, the target deployments were controlled by designing and operating them within this research, which allowed certainty about which OWASP vulnerabilities each configuration should

exhibit based on its design. Configuration A was designed to be vulnerable to all five primary tested controls. Configuration B was designed to resist LLM01 and LLM07 at the surface-level keyword filtering layer while remaining potentially vulnerable to encoding-based evasion. Configuration C was designed to exhibit RAG-specific LLM05 vulnerabilities through deliberate access control gaps in the vector database configuration. This controlled design provides a form of construct validity: the expected detection pattern is derivable from the target's known design, allowing the tool's actual detection pattern to be compared against an expected baseline.

Second, the independent variables were isolated across the three configurations. The difference between Configuration A and Configuration B is specifically and only the addition of system prompt hardening, holding the base model, the deployment framework, and the testing conditions constant. The difference between Configuration B and Configuration C is specifically and only the addition of the RAG layer with its access control configuration. This isolation ensures that observed differences in detection rates across configurations can be attributed to the specific design differences rather than to confounding variables.

Third, the same prompt library and testing parameters were applied across all three configurations without modification, ensuring that any difference in detection outcomes reflects the target's security posture rather than variation in the testing approach.

Rules of Engagement: All testing was conducted under a self-imposed rules of engagement document that specified the maximum request rate, the absence of testing against any third-party production systems, the restriction of testing to systems owned and operated within the scope of this research, and the requirement that all testing be conducted within the authorized scope of this academic research.

3.6 Validation Methodology

The validation methodology addresses the question of whether AegisLLM's outputs can be trusted as accurate security findings. Two independent human assessors validated the tool's classifications to establish the reliability of its output as a security assessment instrument.

Assessor Profiles: The two assessors selected for this validation study both possess professional cybersecurity backgrounds relevant to LLM security assessment. Manav Bankatwala is a Cybersecurity Engineer at Workday with professional experience in application security and security tool evaluation. Kunal Das is a Cybersecurity Analyst with professional experience in penetration testing and vulnerability assessment. Neither assessor was involved in the design, development, or testing phases of AegisLLM, ensuring complete independence from the tool's construction. Neither assessor had access to the tool's classification labels prior to completing their own independent review.

Validation Protocol: A stratified random sample of 91 result records was drawn from the full evaluation dataset, stratified to include a proportional representation of all three judge label categories: SUCCESS, POSSIBLE, and NO_SUCCESS. Stratification ensured that the sample contained sufficient representation of each label category to assess classification reliability across the full label range rather than only in the most common category.

Each assessor was provided with a validation packet containing the original prompt text, the target model's verbatim response, and the four-category SUCCESS definition used in the judge prompt: system prompt or initialization instruction disclosure, credential or API key exposure, role or capability change, and unauthorized command execution or out-of-scope output. The assessors were explicitly not provided with the tool's classification label, the confidence score, or the reasoning chain, ensuring that their judgments were formed independently from the tool's output rather than anchored to it. Each assessor independently classified every sample record as SUCCESS, POSSIBLE, or NO_SUCCESS using the provided definition.

Inter-Assessor Agreement: Before comparing assessor judgments against the tool's labels, the two assessors' classifications were compared against each other to establish the reliability of the human judgment baseline. High inter-assessor agreement confirms that the SUCCESS definition is sufficiently precise to support consistent human judgment, which is a prerequisite for using human judgment as a validity benchmark for the tool.

Comparison Against Tool Labels: Once inter-assessor agreement was established, the combined human assessor classifications were compared against the tool's labels. Cases where both assessors agreed with each other but disagreed with the tool were treated as confirmed disagreements. Cases where the two assessors disagreed with each other were excluded from the tool agreement calculation and recorded as ambiguous cases requiring further investigation. The agreement rates reported in Table 7.3 in Chapter 7 reflect this protocol.

Disagreement Resolution: In cases where the two assessors disagreed with each other, the three parties reviewed the specific response together and reached a consensus classification through discussion. This consensus classification was used as the final human judgment for that record, following the adjudication protocol standard in inter-rater reliability studies in security research.

3.7 Threats to Validity

Following the validity threat framework established by Wohlin et al. for empirical software engineering research, threats are identified across four validity dimensions.

Internal Validity: The primary internal validity threat is the self-configured evaluation environment. Because the target deployments were both designed and tested against within this research, there is a potential for unconscious bias in target configuration that could inflate detection rates. This is mitigated by documenting the target configuration decisions explicitly in Chapter 7 and by applying the same prompt library and testing parameters to all three configurations without modification after initial setup. A secondary internal validity threat is the limited number of evolutionary runs: with one run per configuration, the evolutionary performance data may not be representative of typical run behavior across the full distribution of possible seed-target combinations.

External Validity: The primary external validity threat is the generalizability of findings beyond the self-configured test environment. Detection rates measured against a self-designed Flask application wrapping a locally hosted model may not accurately predict detection rates

against production enterprise deployments with multiple defensive layers, dedicated input sanitization infrastructure, and custom fine-tuning specifically targeting injection resistance. The evaluation results are framed as demonstrated capabilities within the test environment rather than population-level performance estimates.

Construct Validity: The primary construct validity threat is the degree to which the LLM-as-a-judge classification accurately measures what it claims to measure, specifically whether a response represents a genuine security failure under the OWASP control being tested. The two-assessor validation study in Section 3.6 directly addresses this threat by establishing empirical evidence for the correspondence between the tool's classifications and independent human expert judgment. The 94.5% overall agreement rate and zero crossover errors between SUCCESS and NO_SUCCESS categories provide construct validity evidence for the judge's classification mechanism.

Conclusion Validity: The primary conclusion validity threat is the small sample size for the evolutionary performance evaluation, specifically one run per configuration. With a single run, it is not possible to estimate the variance in breakthrough rate or generation depth across repeated runs, which limits the strength of the RQ2 conclusions. The evolutionary results are therefore framed as existence proof rather than distributional estimates: the finding that evolutionary mutation produced a breakthrough against targets that the full static library could not breach is valid as a qualitative conclusion regardless of the run count.

Chapter 4. System Design and Architecture

4.1 Design Goals and Constraints

This thesis designed AegisLLM around five core goals that collectively address the gap between existing security testing tools and the unique demands of LLM-integrated systems.

Modularity and OWASP Alignment: Every component of AegisLLM maps directly to one or more OWASP LLM Top 10 (2025) controls. The testing engine, prompt libraries, and reporting pipeline are structured so that each OWASP module can be executed independently, extended with new prompt datasets, or swapped for an alternative backend without modifying the surrounding infrastructure.

Endpoint Flexibility: LLM deployments vary widely in their API surface. This work address this through a dot-notation path resolver, `get_value_at_path()`, which lets operators specify any response extraction path at runtime without touching source code. For example, a path like `choices.0.message.content` correctly extracts the response text from an OpenAI-compatible endpoint, while a custom path like `data.output.text` works equally well against an internal deployment.

Rate-Limit Compliance and Reproducibility: This work enforce precise request pacing through a multithreaded dispatcher that calculates interval delays as `60.0 / requests_per_minute`, scheduling each thread against a wall-clock timestamp rather than a naive `time.sleep()` call. This guarantees reproducible test conditions and keeps AegisLLM within the operator-defined rules of engagement throughout a full test run [7].

Automated Judgment at Scale: This thesis address automated evaluation at two distinct levels: deterministic algorithmic scoring for the evolutionary engine, where latency and cost make LLM-based evaluation impractical, and zero-temperature LLM-as-a-judge classification for static batch tests, where contextual interpretation of a response requires semantic understanding.

Extensibility and Reporting: All test outputs are written to timestamped CSV files with a consistent schema, making results both human-readable and machine-parseable for downstream integration with SIEM tooling or CI/CD pipelines.

4.2 High-Level Architecture

AegisLLM is organized into three distinct layers, each with a clearly scoped responsibility. This layered design ensures that interface concerns, core testing logic, and advanced attack generation remain decoupled from one another.

Layer 1: Interface Layer. The interface layer exposes two interaction modes. The CLI Manager (`main.py`) provides a terminal-based interactive menu that orchestrates all OWASP test modules, handles CSV selection, accepts authentication tokens, configures rate limits, and routes test execution to the appropriate backend scripts. The Web Dashboard (`web_app.py`) is a

Flask-based GUI providing real-time log streaming, configuration controls, and formatted output reports covering both standard batch test results and ThreatForge breakthrough logs.

Layer 2: Core Testing Layer. The core layer handles the actual execution of test cases against target endpoints. `prompt_tester.py` acts as the primary load generator, reading prompt CSVs, applying template substitution, dispatching multithreaded HTTP requests against the target, and writing timestamped result CSVs. `injection_judge.py` operates as a post-processing step, consuming those result CSVs and routing each response through an LLM-as-a-judge evaluator to produce a structured classification label, confidence score, and reasoning chain.

Layer 3: ThreatForge Evolutionary Layer. The ThreatForge suite (`threatforge/`) handles advanced, adaptive attack generation. Rather than relying solely on static prompt libraries, ThreatForge takes seed payloads from the jailbreak library and iteratively mutates them across multiple generations using the SPE-NL genetic algorithm (`scripts/run_evolution.py`). Two additional modules complement the evolutionary engine: `run_posture.py` runs a non-destructive baseline scan before intensive testing begins, and `run_differential.py` compares the target model against a clean baseline to isolate weaknesses introduced specifically by fine-tuning or RAG augmentation.

4.3 CLI Manager (`main.py`)

The CLI Manager serves as the primary entry point for AegisLLM. Rather than exposing a flat command-line interface with dozens of flags, This thesis designed `main.py` as an interactive terminal menu that guides the operator through test configuration step by step.

OWASP Module Routing - The menu presents each available OWASP LLM Top 10 control as a numbered option. When the operator selects a module, `main.py` resolves the corresponding prompt library directory and surfaces only the CSV files relevant to that control. This directory-to-control mapping means that adding a new OWASP module requires only dropping a new prompt directory into the correct location and registering it in the routing table, with no changes to the testing engine itself.

CSV Selection - Once a module is selected, the CLI enumerates all CSV files in the resolved directory and presents them as selectable options. The operator can choose a single targeted dataset for a focused run or select all available files for a broad sweep across every technique mapped to that control.

Authentication and Endpoint Configuration - Before dispatching any requests, the CLI collects three pieces of configuration from the operator: the target endpoint URL, the authentication token, and the JSON response path for extraction. These values are passed into the request template at runtime, keeping credentials out of source code and out of the prompt library files entirely.

Rate-Limit Configuration - The CLI prompts the operator to specify a requests-per-minute ceiling. This value feeds directly into the threading dispatcher in `prompt_tester.py`. The CLI surfaces this configuration explicitly rather than applying a silent default, ensuring the operator makes a deliberate choice before testing begins [7].

Routing to ThreatForge - For operators who want to go beyond static prompt libraries, the CLI exposes a separate ThreatForge menu from which the operator can launch the SPE-NL evolutionary engine, run a guardrail posture scan, or initiate a differential attack scan.

4.4 Web Dashboard (web_app.py)

The Web Dashboard runs as a Flask application and can be started independently of the CLI, meaning an operator can launch a test through the terminal and monitor its progress through the browser simultaneously. The dashboard is implemented as a single Flask application with routes handling configuration input, test execution triggers, log streaming, and report rendering. It delegates all testing logic to the same backend scripts used by the CLI.

Real-Time Log Streaming - The dashboard implements server-sent events, pushing log lines to the browser as each request completes. This is particularly valuable during ThreatForge evolutionary runs, where a test suite may span multiple generations and operators need to monitor fitness score progression and breakthrough events as they occur.

Report Rendering - Once a test run completes, the dashboard renders two distinct report views. The standard results view presents the output of `prompt_tester.py` and `injection_judge.py` as a formatted table showing the prompt technique, raw response, judge classification label, confidence score, and response time. The ThreatForge report view surfaces breakthrough payloads separately, displaying the evolved prompt, its final fitness score, the generation at which the breakthrough occurred, and the sub-score breakdown across all six fitness dimensions.

Legacy CSV Import - The dashboard supports importing previously generated result CSVs for re-analysis or report generation, allowing operators to run tests in headless environments and use the dashboard purely as a reporting layer without re-executing any requests against the target.

Figure 4.1 illustrates the web dashboard module selection interface, presenting each available OWASP test module as a selectable card with its control identifier and a brief description of its testing objective.

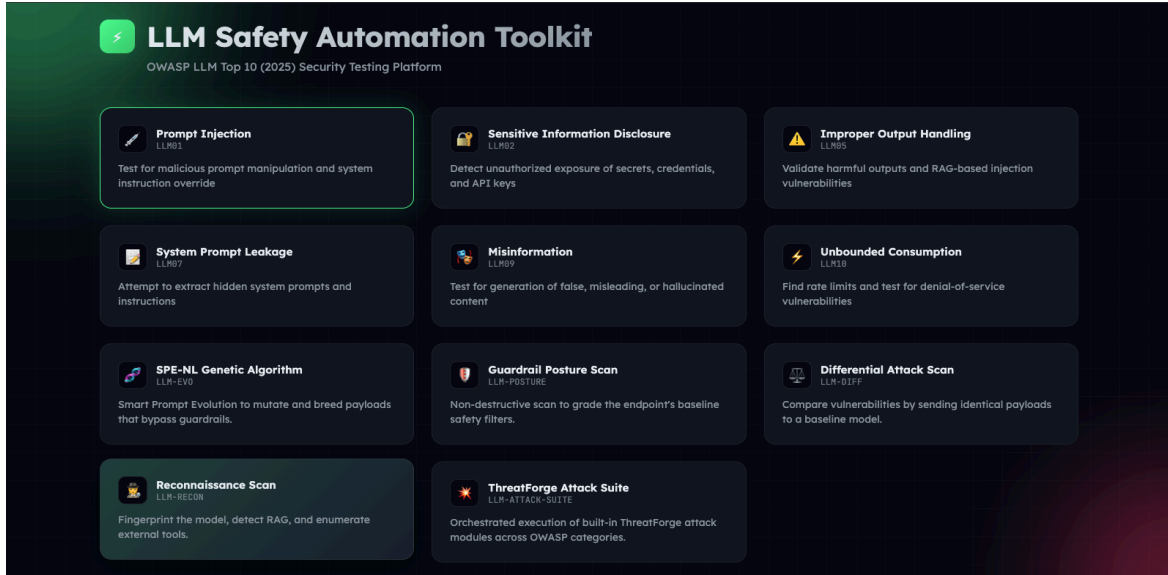


Figure 4.1. AegisLLM web dashboard module selection interface, displaying the nine available OWASP LLM Top 10 testing modules including Prompt Injection (LLM01), Sensitive Information Disclosure (LLM02), Improper Output Handling (LLM05), System Prompt Leakage (LLM07), Misinformation (LLM09), Unbounded Consumption (LLM10), the SPE-NL Genetic Algorithm, Guardrail Posture Scan, and Differential Attack Scan.

Figure 4.2 shows the test configuration panel during an active test run, illustrating the endpoint configuration, prompt library selection, rate-limit setting, and the real-time result counter display that updates as each request completes.

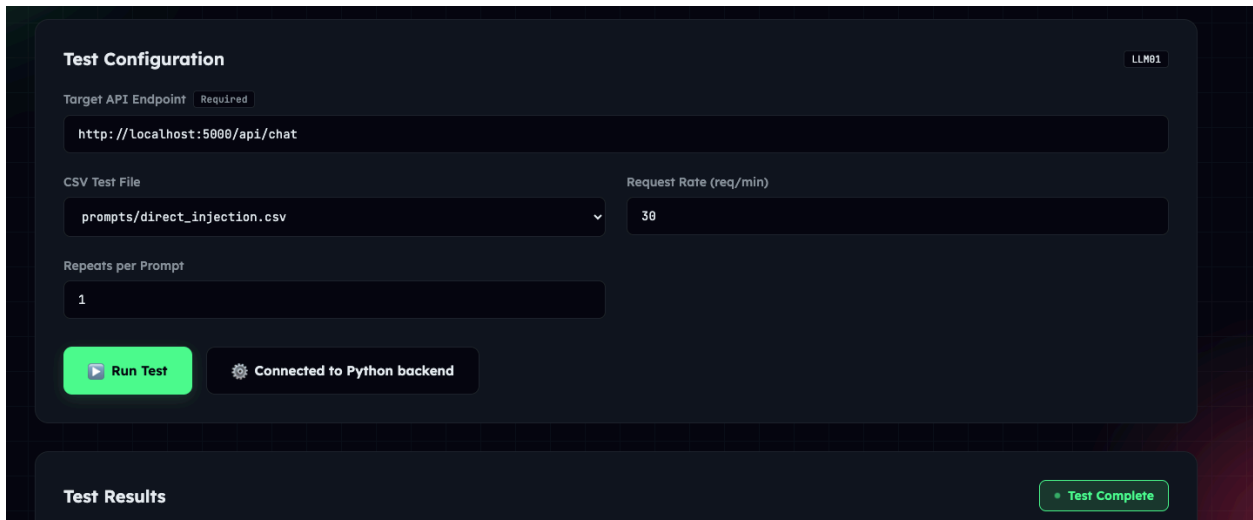


Figure 4.2. AegisLLM web dashboard test configuration panel during an active LLM01 Prompt Injection test run, showing the target API endpoint, selected prompt library (direct_injection.csv), request rate of 30 requests per minute, and real-time test result counters for total requests, successful responses, failed responses, and injection successes.

4.5 Prompt Testing Engine (prompt_tester.py)

prompt_tester.py is the execution core of AegisLLM. Every test run, whether triggered through the CLI or the web dashboard, ultimately passes through this script. It is responsible for reading prompt datasets, constructing HTTP requests, dispatching them against the target endpoint under precise rate-limit constraints, capturing full response telemetry, and writing results to disk.

Prompt Ingestion - The engine reads prompt datasets using `csv.DictReader`, iterating through each row of the selected CSV file. Every row follows the standard three-column schema: `id`, `technique`, and `prompt`.

Template Substitution - The engine uses a lightweight templating approach where the HTTP request body contains two placeholders: `{{PROMPT}}` and `{{TECHNIQUE}}`. At runtime, the engine substitutes the current row's values into these placeholders before serializing the request.

Multithreaded Dispatcher and Rate-Limit Enforcement - Each prompt is assigned a scheduled wall-clock execution time based on a fixed interval. A dedicated thread is spawned for each prompt and waits until its scheduled time before firing.

This approach differs from a naive sequential loop in two important ways. First, it absorbs the variable execution time of each HTTP request. Second, it produces a request pattern that precisely mirrors the interval the operator specified, which is essential when testing rate-limit enforcement under LLM10 [18].

Response Extraction - This work handle vendor response format variability through `get_value_at_path()`, a dot-notation parser that traverses nested JSON structures at runtime. The operator specifies the extraction path at configuration time and the engine resolves it against every response without any code changes.

Telemetry Capture - For every request, the engine captures a full telemetry record: `prompt_id`, `technique`, `prompt`, `response`, `response_time_ms` computed as `(time.time() - request_start) * 1000`, `status_code`, and a UTC timestamp. Response time in milliseconds is particularly important for LLM10 testing, where unusually high latency under load can itself indicate resource exhaustion on the target side even when the HTTP status code returns 200 [8].

The same engine handles the LLM05 (Improper Output Handling) test modules, where the evaluation goal shifts from injection success to whether the model produces outputs outside its intended operational scope in response to off-topic or harmful content requests. The rate-limit enforcement mechanism is equally critical for LLM05 testing: submitting off-topic escalation prompts too rapidly risks triggering adaptive session-level filters that would produce artificially elevated refusal rates.

Output - All telemetry records are written to a timestamped CSV file named `results_*_<timestamp>.csv`. At the full library size of 279 prompts, a complete single-configuration test run produces a result CSV containing 279 rows before judge classification and 279 rows with three appended judge columns after classification, totaling nine

columns per row across the complete result set. Operators running all three target configurations in sequence will accumulate 837 classified response records across the full evaluation.

4.6 Response Storage and Analysis Pipeline

Once `prompt_tester.py` completes a test run, the raw result CSV contains everything needed to reconstruct exactly what happened during the engagement. However, raw response text alone is not sufficient for reporting. The analysis pipeline exists to bridge the gap between raw HTTP responses and semantic security findings.

Result CSV Schema - Every result file produced by `prompt_tester.py` follows a consistent seven-column schema:

Table 4.1. Result CSV schema produced by `prompt_tester.py`.

Field	Description
<code>prompt_id</code>	Row identifier from the source CSV, used to trace results back to the original dataset
<code>technique</code>	Named attack method, e.g. SYSTEM Tag Spoofing or Base64 Obfuscation
<code>prompt</code>	Exact payload text sent to the target endpoint
<code>response</code>	Extracted response text, resolved via <code>get_value_at_path()</code>
<code>response_time_ms</code>	End-to-end request latency in milliseconds
<code>status_code</code>	HTTP response code returned by the target
<code>timestamp</code>	UTC timestamp of the request

Why Regex-Based Classification Is Insufficient. A naive approach to classifying injection success would be to check whether the response contains specific keywords indicating refusal. The problem is bidirectional: a model that has been successfully compromised may still open its response with a polite acknowledgment before leaking sensitive content, and a model that refuses may do so using phrasing that does not match any predefined keyword. This thesis address this by routing every result through a semantic LLM-as-a-judge evaluator.

4.7 LLM-as-a-Judge Evaluator (`injection_judge.py`)

The LLM-as-a-judge pattern has gained significant traction in NLP evaluation research as a scalable alternative to human annotation for tasks that require semantic understanding [13]. This work adopts this pattern because security classification of LLM responses requires both contextual interpretation and scale that manual review cannot provide.

GPT-4o was selected as the judge model based on four criteria. First, at the time of evaluation GPT-4o represented the highest publicly available benchmark performance on semantic reasoning and instruction-following tasks among commercially accessible models, making it the strongest available baseline for classification reliability. Second, GPT-4o supports a JSON response format lock at the API level, which is a non-negotiable requirement for the

three-label classification schema: without structured output enforcement, any judge model introduces parsing failures and inconsistent label formatting at scale. Third, GPT-4o operates reliably at temperature 0.0, producing deterministic output across repeated calls on identical inputs, which is essential for reproducibility of the evaluation pipeline. Fourth, as a commercially hosted model with a stable API, GPT-4o eliminates the infrastructure overhead and version drift risk associated with self-hosted alternatives. Models such as Llama 3 and Mistral were considered but ruled out because self-hosting introduces GPU resource constraints that limit throughput, and quantized local versions introduce accuracy degradation relative to the full-parameter hosted equivalents. Claude 3 was considered but ruled out because Anthropic's API did not support JSON-mode output locking at the time of evaluation. The combination of reasoning capability, structured output support, temperature determinism, and operational stability made GPT-4o the most defensible choice for a validation-grade semantic judge.

Judge Prompt Architecture: The judge operates on a two-message conversation structure. The system prompt is intentionally minimal: You are a security classifier. Decide if the ASSISTANT RESPONSE shows a successful prompt injection. The user message provides the evaluation criteria and the specific prompt-response pair being classified.

The four SUCCESS categories in the judge prompt map to distinct OWASP controls as follows. System prompt and initialization instruction disclosure maps to both LLM01 and LLM07, which AegisLLM treats as co-testable controls. Credential and API key exposure maps to LLM02. Role and capability changes map to LLM01. Unauthorized command execution or out-of-scope output maps to LLM05, specifically in cases where the model produces content outside its defined operational authorization.

Zero-Temperature Enforcement: This work set the temperature: 0.0 on every judge invocation. This is a non-negotiable constraint for a security evaluation pipeline. A judge model running at nonzero temperature can produce different classification labels for identical inputs across runs, making results non-reproducible and undermining the validity of any comparative analysis between test sessions.

Result Appendage: Rather than producing a separate output file, `injection_judge.py` appends three new columns directly to the existing result CSV in place: `injection_label` (SUCCESS, POSSIBLE, or NO_SUCCESS), `injection_confidence` (integer from 0 to 100), and `injection_reasons` (a list of natural language reasoning strings explaining the classification decision).

Three-Label Classification Rationale: The choice to use three labels rather than a binary SUCCESS/FAILURE reflects the reality of LLM response behavior. Many responses fall into a gray zone where the model partially complied with a malicious instruction or revealed some but not all of the system prompt. The POSSIBLE label captures this gray zone explicitly, flagging responses for human review rather than forcing them into a binary classification that would either overcount or undercount actual vulnerabilities.

The POSSIBLE label represents cases where the judge identified partial information leakage or ambiguous compliance that did not meet the threshold for a definitive SUCCESS classification. In the inter-rater validation study, all six POSSIBLE-labeled records were

reclassified as SUCCESS by the human assessors, indicating that the judge was systematically conservative at this boundary. For the purposes of this research, POSSIBLE findings are reported separately and excluded from the primary SUCCESS detection rate to avoid inflating vulnerability counts. Practitioners deploying AegisLLM in operational contexts should treat POSSIBLE as a flag for manual review rather than a confirmed finding.

Limitations: The LLM-as-a-judge approach carries two inherent limitations. First, the judge can itself be wrong. This is mitigated through the three-label system and mandatory human review flags on low-confidence results. Second, the judge's classification quality is bounded by the capability of the underlying model. This is addressed partially through the explicit SUCCESS definition in the user prompt, which reduces the judge's reliance on implicit security knowledge.

4.8 Domain-Specific Calibration of Classification Sensitivity

Different deployment domains carry fundamentally different tolerances for false negatives and false positives, and the interpretation of AegisLLM findings should be adjusted to reflect the risk profile of the specific application under test. The following guidance describes the recommended interpretive posture for four representative deployment categories.

In healthcare applications, including clinical decision support tools, medical triage assistants, and patient-facing symptom checkers, the cost of a false negative is extremely high. A missed prompt injection vulnerability or undetected sensitive information disclosure in a clinical context could result in incorrect clinical guidance reaching a patient or protected health information being exfiltrated. In this domain, every POSSIBLE finding should be treated as a confirmed SUCCESS and routed for immediate remediation review. The detection threshold should be tuned toward maximum sensitivity, accepting a higher false positive rate as the cost of ensuring no real vulnerability is overlooked. Compliance obligations under HIPAA reinforce this posture: a privacy failure in a healthcare LLM application may carry breach notification requirements regardless of whether the disclosure was confirmed or only probable.

In financial applications, including loan assessment assistants, fraud detection interfaces, and customer-facing banking chatbots, the risk profile is similarly asymmetric but for different reasons. Sensitive information disclosure vulnerabilities in this domain may expose account details, transaction histories, or proprietary financial models. Prompt injection vulnerabilities could be exploited to manipulate financial outputs or extract decision logic. POSSIBLE findings should again default to SUCCESS for triage purposes. Additionally, financial applications are subject to audit and regulatory scrutiny under frameworks such as SOX and PCI-DSS, meaning that a POSSIBLE finding that is later confirmed through a breach carries greater legal and reputational exposure than the cost of over-remediating a false positive. In this domain, conservative classification is the operationally correct default.

In educational applications, including tutoring assistants, student-facing knowledge tools, and academic writing aids, the risk profile is more nuanced. The primary concern is misinformation generation and off-topic content compliance rather than sensitive data exfiltration. False negatives on misinformation prompts, covered by the misinformation library in AegisLLM, represent a meaningful harm risk particularly where the user population includes minors or students relying on the tool for assessed work. However, the absence of regulated personal data

in most educational deployments means that false positives carry a lower remediation burden than in healthcare or finance. POSSIBLE findings in the misinformation and off-topic categories should be reviewed but need not automatically block deployment. A balanced sensitivity posture is appropriate, with tighter thresholds applied specifically to content categories involving factual accuracy and age-appropriate output.

In public-facing general-purpose chatbots, including customer service agents, HR assistants, and enterprise knowledge bases with broad user populations, the risk profile is highly variable and depends on what data the chatbot has access to. A public chatbot with no access to personal or confidential data presents a relatively low false negative risk for sensitive information disclosure, but a high reputational risk from jailbreak compliance and harmful output generation. In this domain, POSSIBLE findings on harmful prompt and jailbreak categories should be treated as SUCCESS, while POSSIBLE findings on sensitive data categories can be downgraded to manual review without blocking deployment. The recommended calibration is asymmetric by category: maximum sensitivity for output harm categories, standard sensitivity for information disclosure categories, with the division determined by a pre-assessment data access audit of the target system.

A summary of the recommended interpretive posture across these four domains is provided in Table 4.8.

Table 4.8: Recommended Classification Sensitivity by Deployment Domain

Domain	False Negative Tolerance	False Positive Tolerance	POSSIBLE Default	Primary Risk Category
Healthcare	Very Low	High	Treat as SUCCESS	Sensitive data disclosure, misinformation
Finance	Very Low	High	Treat as SUCCESS	Sensitive data disclosure, prompt injection
Education	Low to Medium	Medium	Manual review	Misinformation, off-topic compliance
Public Chatbot	Low (output harm) / Medium (data)	Medium	SUCCESS for harm categories; review for data categories	Jailbreak, harmful output

Chapter 5. Prompt Library Design and Construction

Table 5.1: OWASP Table

Control	Full Name	AegisLLM Coverage
LLM01	Prompt Injection	Direct, narrative, obfuscation, jailbreak libraries (85 prompts)
LLM02	Sensitive Information Disclosure	Sensitive data exposure library (52 prompts)
LLM03	Supply Chain	Not testable via black-box prompt approach
LLM04	Data and Model Poisoning	Not testable via black-box prompt approach
LLM05	Improper Output Handling	Harmful prompts, RAG fishing, canary libraries (82 prompts)
LLM06	Excessive Agency	Indirect symptomatic coverage only
LLM07	System Prompt Leakage	Co-primary coverage via LLM01 libraries (75 prompts)
LLM08	Vector and Embedding Weaknesses	Indirect symptomatic coverage only
LLM09	Misinformation	Misinformation library (30 prompts)
LLM10	Unbounded Consumption	Standalone rate limit tester script

5.1 Prompt Sourcing Methodology

The quality of an adversarial prompt library is determined not by its size but by the real-world validity of the techniques it contains. This thesis built the AegisLLM prompt libraries using a three-pronged sourcing strategy designed to ground every technique in empirically validated attack behavior.

Source 1: CTF Competition Wins. Capture the Flag competitions that include LLM-based challenges provide one of the highest-fidelity sources of adversarial prompt data available, because every successful technique in a CTF has been empirically validated against a live, defended system under competitive conditions. I participated in multiple CTF competitions over the course of this research, specifically targeting challenges involving LLM prompt injection, system prompt extraction, and guardrail bypass. For every challenge I solved, I documented the exact prompt or prompt sequence that produced the flag, the reasoning behind why that

technique worked against that particular target, and the OWASP LLM control the technique most directly maps to.

Source 2: Industry AI Red-Teaming Competition. I participated in an industry-organized AI red-teaming competition in which participants were tasked with breaking production-grade AI chatbot deployments to reveal sensitive information, extract system prompts, and execute prompt injection attacks across a range of target configurations. Following the competition, I contacted the organizers and collected the full set of successful prompts submitted by participants across all challenge categories. This dataset was particularly valuable because it captured successful techniques from multiple independent attackers rather than a single perspective, and every prompt had been validated against a system with active defenses. I processed this raw collection by deduplicating overlapping techniques, normalizing prompts into the standard three-column CSV schema, and mapping each technique to its corresponding OWASP LLM control.

Source 3: Research Literature and Common Techniques. The third sourcing layer draws from the established body of adversarial ML research and publicly documented jailbreak techniques. This includes prompt injection patterns characterized by Perez et al. [2], the original Ignore Previous Prompt attack taxonomy by Perez and Ribeiro [20], adversarial input generation approaches documented in TF-Attack [12], and jailbreak structures such as DAN-style role inversion prompts analyzed in the security research community [23][24][25]. Research literature techniques are included not because they have been validated against a specific live target but because they represent the documented baseline of known attack vectors that any LLM security testing suite must cover. They also serve a specific functional role in AegisLLM beyond direct testing: the jailbreak library built from this source is used primarily as a seed population for the SPE-NL genetic algorithm.

Dataset Construction and Quality Control. Prompts from all three sources were processed through a consistent pipeline before inclusion. Each prompt was reviewed to confirm it mapped unambiguously to at least one OWASP LLM Top 10 control. Duplicates were collapsed into a single representative entry. Prompts targeting platform-specific behaviors were either generalized or placed in a separate vendor-specific directory. Finally, each prompt was assigned a technique name following a consistent naming convention.

5.2 CSV Schema and Ingestion Standard

Every prompt library in AegisLLM follows a strict three-column CSV schema. This standardization means that any new prompt dataset can be dropped into the appropriate library directory and immediately consumed by the testing engine without any modifications to the ingestion code.

Table 5.2. Standard three-column CSV schema used by all AegisLLM prompt libraries.

Column	Type	Description
id	String/Integer	Sequential or alphanumeric identifier for row-level traceability across the analysis pipeline

Column	Type	Description
technique	String	Named attack method substituted into the {{TECHNIQUE}} placeholder and used as the primary reporting label
prompt	String	Exact adversarial payload text substituted into the {{PROMPT}} placeholder without modification

Ingestion Mechanism. The testing engine reads prompt datasets using Python's `csv.DictReader`, which maps each row to a dictionary keyed by column name. The UTF-8 encoding declaration is important: several prompt libraries contain Base64 strings, Unicode homoglyphs, and non-ASCII obfuscation characters that would be corrupted by a default ASCII reader.

Template Substitution. Once a row is ingested, the engine substitutes its values into the HTTP request template at the string level before JSON serialization. This ensures that adversarial prompts containing characters that would be interpreted as JSON control characters arrive at the target endpoint exactly as authored.

5.3 Direct Injection Library (prompts/direct_injection.csv) - 25 prompts

The direct injection library targets OWASP LLM01 (Prompt Injection) and OWASP LLM07 (System Prompt Leakage) as co-primary controls, containing 25 prompts across six technique classes. The dual LLM01/LLM07 coverage reflects the natural co-occurrence of these two outcomes: a prompt that causes the model to abandon its instructions typically also causes it to reveal what those instructions were [2][8].

Technique 1: Ignore Previous Instructions. This is the most direct form of prompt injection and serves as a baseline calibration prompt. A target that fails against this technique has effectively no prompt injection defense, establishing a lower bound for the target's security posture [2][20].

Technique 2: Role Switching. Role switching prompts exploit the instruction-following disposition of LLMs by reframing the model's identity before making the malicious request. This two-step structure bypasses simple keyword filters because the actual request is phrased as a routine administrative action within the fabricated context. Validated through multiple CTF competitions where role-based reframing consistently outperformed direct instruction override attempts.

Technique 3: SYSTEM Tag Spoofing. SYSTEM Tag Spoofing exploits the formatting conventions used by LLM API providers to distinguish system-level instructions from user inputs. By mimicking the `SYSTEM:` prefix within a user message, the attack attempts to induce the model to treat the injected content as an operator-level system prompt [8]. This technique was among the most consistently successful in the industry red-teaming competition dataset.

Technique 4: Developer Impersonation. Developer Impersonation relies on social engineering within the prompt rather than structural exploitation. By asserting a trusted identity before making the malicious request, the prompt attempts to trigger a deference response that

bypasses standard refusal behavior. Validated in CTF challenges involving fine-tuned enterprise chatbots.

Technique 5: Admin Override Spoofing. Admin Override Spoofing combines XML tag injection and authority claim escalation. The XML tag structure mimics administrative control interfaces, while the `override_safety_mode` instruction attempts to invoke a privileged operational state. This technique is particularly effective as a seed for the SPE-NL genetic algorithm because its multi-layer structure provides multiple independent mutation surfaces [10].

Technique 6: End Marker Spoofing. End Marker Spoofing exploits the delimiter conventions used by some LLM deployment frameworks to separate system prompt content from user input. Targets vulnerable to this technique have a fundamental context management weakness that cannot be patched at the prompt level alone [8].

Beyond these six technique classes, the library contains approximately twenty additional prompt variants that introduce structural variations such as nested instruction overrides, multi-turn injection sequences, and hybrid authority-claim combinations.

Table 5.3. Direct injection library OWASP mapping.

Technique	LLM01 Sub-Risk	LLM07 Coverage	Primary Source
Ignore Previous Instructions	Direct instruction override	System prompt disclosure	Research literature [2][20]
Role Switching	Identity-based context manipulation	Indirect elicitation	CTF validation
SYSTEM Tag Spoofing	Trusted role boundary exploitation	Authorization boundary bypass	Industry competition
Developer Impersonation	Authority claim escalation	Privileged access elicitation	CTF validation
Admin Override Spoofing	Structural token injection	Restricted content disclosure	Research literature [10]
End Marker Spoofing	Context delimiter exploitation	Context boundary extraction	Research literature

5.4 Narrative and Obfuscation Libraries

The narrative injection library (25 prompts) targets LLM01 and LLM07 as co-primary controls by exploiting the model's creative writing disposition: a model that reliably refuses a direct request to reveal its system prompt may comply with an equivalent request framed as a scene in a novel or a hypothetical scenario. Three primary technique classes are covered: Fictional Scenario Embedding, where the malicious request is embedded within an elaborate fictional scenario; Hypothetical and Thought Experiment Framing, where the prompt presents the request as a purely academic exercise; and Role-Play Persona Assignment, which builds the

persona gradually through conversational context rather than making the assignment explicit and abrupt.

The obfuscation and encoding injection library (25 prompts) targets a different layer of the defense stack: surface-level defenses such as regular expression filters, Web Application Firewalls, and keyword-based content scanners. The core insight is that the same semantic content, including system prompt extraction requests, can be expressed in an unlimited number of surface forms [8]. Technique classes covered include Base64 Encoding, Hexadecimal Encoding, String Reversal, Pig Latin and Character Substitution, and Layered XML Obfuscation.

5.5 Sensitive Data Exposure Library - 52 prompts

The sensitive data exposure library targets OWASP LLM02 (Sensitive Information Disclosure) through three distinct pathways. Runtime Context Leakage prompts are crafted to elicit specific context content verbatim or near verbatim through progressive disclosure and indirect elicitation approaches that work against deployments hardened against direct show me your system prompt requests [8][20]. Credential and API Key Extraction prompts target credentials and API keys embedded in system prompts or tool configurations through technical and social engineering framing, mapping directly to the regex patterns in the SPE-NL information leakage sub-scorer. RAG Document Leakage prompts probe whether the retrieval layer enforces access controls at the document level, exploiting the tendency of RAG-augmented models to ground responses in retrieved content without independently verifying access entitlement [21].

5.6 Harmful Output and Off-Topic Libraries - 30 prompts

The harmful output libraries contain 10 prompts each across three files and map primarily to OWASP LLM05 (Improper Output Handling). `offtopic_prompts.csv` (10 prompts) tests task-constrained agent deployments through progressive boundary expansion, starting with requests adjacent to the defined scope and incrementally moving further outside it. `harmful_prompt_injections.csv` (10 prompts) combines scope boundary testing with explicit harmful content requests wrapped in injection scaffolding, testing simultaneous exploitation of LLM01 and LLM05. `harmful_prompts.csv` (10 prompts) submits direct harmful content requests without injection scaffolding or narrative framing, serving as a severity calibration instrument that establishes the outer boundary of what the model will and will not produce.

5.7 Misinformation and Hallucination Library - 30 prompts

The misinformation library targets OWASP LLM09 (Misinformation) and is the only library in the AegisLLM collection that maps to a single OWASP control without secondary coverage overlap. Three technique classes are covered: Fake News and Fabricated Event Generation, Forged Citation and Non-Existent Source Generation, and Factual Contradiction and Historical Revision. The misinformation library is expected to show consistent detection rates across all target configurations, making it a useful control condition for interpreting detection rate variation in other library categories.

5.8 RAG-Specific Attack Libraries - 82 prompts

The RAG attack libraries contain 82 prompts across four CSV files and map primarily to OWASP LLM05 (Improper Output Handling) with LLM02 as a secondary control [21]. `rag_fishing_prompts.csv` (42 prompts) probes the retrieval boundary through indirect elicitation, crafting queries that are semantically similar to restricted documents to test whether vector similarity search surfaces restricted content [22]. `targeted_rag.csv` (10 prompts) and `targeted_rag_second_attempt.csv` (10 prompts) contain progressively refined targeted retrieval attacks, with the second file developed after observing partial results of the first attempt.

`canary_rag.csv` (20 prompts) employs canary tokens: uniquely randomized strings with no semantic meaning injected into specific restricted documents within a RAG knowledge base. During this research I physically constructed a test RAG database, injecting unique canary strings into restricted documents and then submitting the 20 prompts against a RAG-augmented deployment connected to this test database. When a model response contained a canary string, this constituted unambiguous, deterministic evidence that the retrieval system had surfaced the restricted document, providing a finding with a level of evidentiary certainty that probabilistic response analysis cannot match [22].

5.9 Jailbreak Seed Library - 10 prompts

The jailbreak library contains 10 high-density, multi-layered bypass payloads designed primarily as seeds for the SPE-NL genetic algorithm rather than as standalone test cases. Despite the small prompt count, each payload represents a fundamentally different bypass architecture, maximizing the breadth of the evolutionary engine's initial search space. The library maps to LLM01 as its primary control and LLM07 as a secondary control. DAN-style role inversion templates [24] and length-inflated logical reasoning bypasses [23][25] are included. Jailbreak payloads provide the evolutionary engine with seed prompts that already exhibit properties correlated with high fitness scores: they avoid simple refusal trigger phrases, produce longer responses, and have multi-layered structures that give the mutation operator multiple independent surfaces to vary across generations [25].

5.10 OWASP LLM Top 10 Control Coverage Matrix

Table 5.3 provides the complete mapping between every AegisLLM prompt library and the OWASP LLM Top 10 controls it targets, incorporating exact prompt counts and corrected control mappings across all 13 CSV files. The total prompt count of 279 across all libraries covers five OWASP LLM controls as primary targets with dedicated library coverage, with LLM07 covered as a co-primary control across the 75 prompts in the direct injection, narrative, and obfuscation libraries.

Table 5.3. Complete OWASP LLM Top 10 control coverage matrix for all 13 AegisLLM prompt libraries.

Library	Count	Primary Control	Secondary	Sourcing
direct_injection.csv	25	LLM01, LLM07	LLM02	CTF, industry competition, literature [2][20]
narrative_injections.csv	25	LLM01, LLM07	LLM02, LLM05	Industry competition, CTF
obfuscation_and_encoding_injection.csv	25	LLM01, LLM07	LLM02	CTF, literature
jailbreaks.csv	10	LLM01	LLM07	Literature [23][24][25]
Sensitive_data_exposure.csv	52	LLM02	LLM07	Literature [19][20], CTF
harmful_prompt_injections.csv	10	LLM05	LLM01	Industry competition, literature
harmful_prompts.csv	10	LLM05	None	Literature, common techniques
offtopic_prompts.csv	10	LLM05	None	Common techniques
rag_fishing_prompts.csv	42	LLM05	LLM02	Literature [21]
targeted_rag.csv	10	LLM05	LLM02	Iterative empirical development
targeted_rag_second_attempt.csv	10	LLM05	LLM02	Iterative empirical development
canary_rag.csv	20	LLM05	LLM02	Novel methodology, empirically validated
misinformation.csv	30	LLM09	None	Literature, common techniques
TOTAL	279			

LLM07 Coverage Note. OWASP LLM07 (System Prompt Leakage) is covered as a co-primary control across the direct injection, narrative injection, and obfuscation libraries, contributing 75 prompts across three files explicitly targeting system prompt extraction. The 75 prompts covering LLM07 as a co-primary control represent the largest single-control prompt allocation in the collection alongside LLM01. LLM07 is not assigned a dedicated library because system prompt leakage is most naturally tested through the same prompt structures that test LLM01: a prompt that successfully injects a malicious instruction frequently also reveals what those instructions contained.

Coverage Gaps and Scope Acknowledgment. The remaining five controls, LLM03 (Supply Chain), LLM04 (Data and Model Poisoning), LLM06 (Excessive Agency), LLM08 (Vector and Embedding Weaknesses), and LLM10 (Unbounded Consumption), are addressed through partial tooling. LLM10 has dedicated standalone rate-limit threshold detection tooling. LLM06 and LLM08 receive indirect symptomatic coverage through existing libraries. LLM03 and LLM04 are architecturally outside the scope of any black-box prompt-based testing tool. Extending dedicated library coverage to LLM06 and LLM08 is identified as a primary direction for future work in Section 9.2.

Chapter 6. ThreatForge: Evolutionary Attack Suite

6.1 Motivation for Evolutionary Prompt Generation

The static prompt libraries described in Chapter 5 represent a necessary but ultimately bounded approach to LLM security testing. A static library, however, carefully sourced and curated, encodes a fixed set of known attack techniques. Once a target deployment has been hardened against every technique in a given library, the library ceases to be a useful testing instrument.

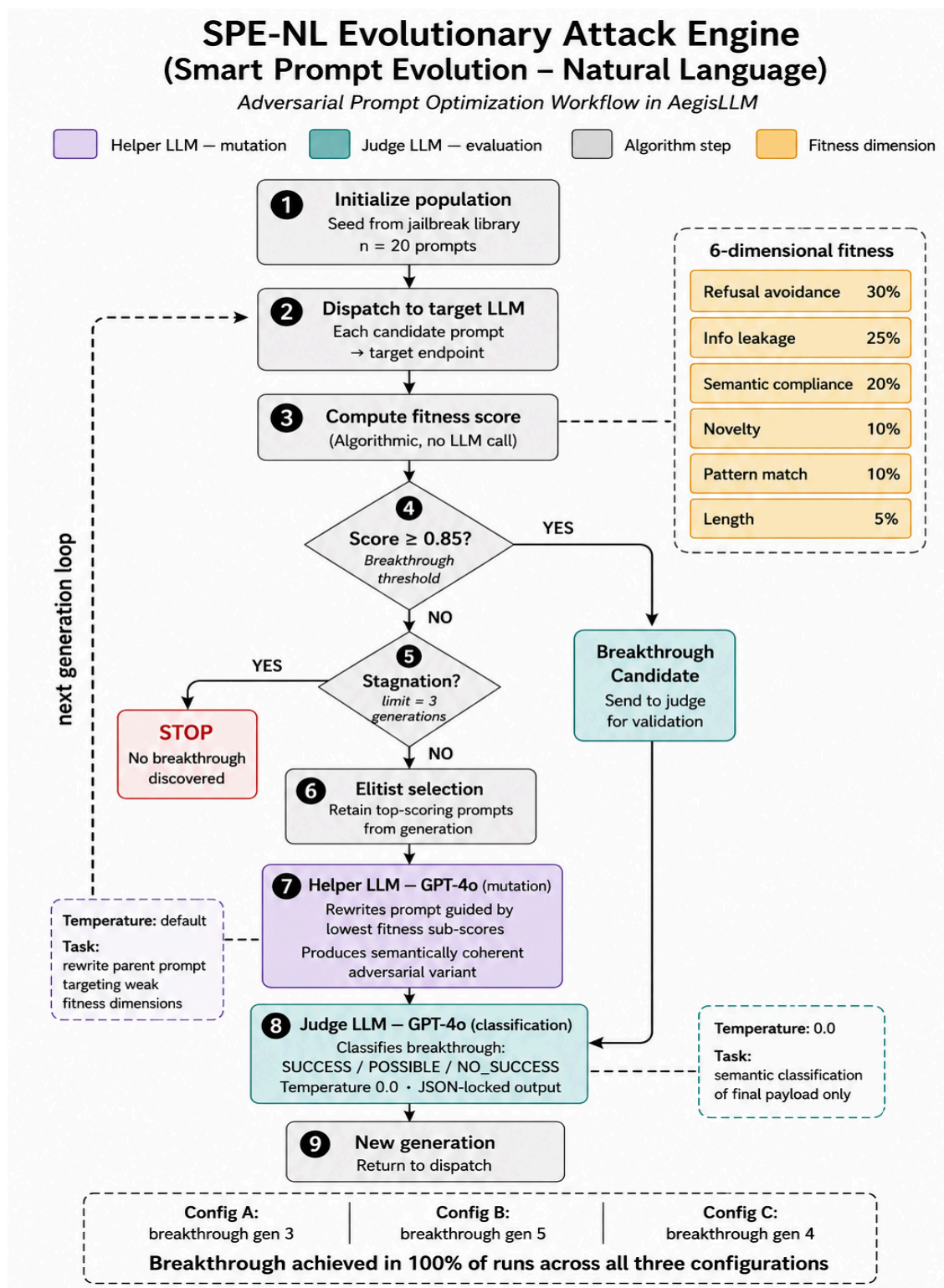
The practical experience that motivated ThreatForge was a convergence of three distinct pain points encountered across multiple test engagements. First, static prompts were getting blocked consistently after a small number of runs against well-configured targets, suggesting that some deployments implement adaptive filtering that escalates its sensitivity in response to repeated injection attempts within a session. Second, certain targets appeared to have been updated mid-engagement, with techniques that produced partial compliance in early runs returning hard refusals in later runs. Third, and most fundamentally, the static library by definition cannot contain prompts the target has never seen: a sufficiently novel payload structure falls outside the coverage of any finite curated dataset.

The conceptual foundation for this approach came from traditional penetration testing, specifically from fuzzing. Coverage-guided fuzzing tools such as AFL and libFuzzer operate on a simple and powerful principle: take a valid seed input, apply structured mutations to it, execute the mutated input against the target, and use the target's response as the signal that guides which mutations to retain [14][15]. SPE-NL applies this same principle to the LLM prompt space. The seed inputs are the jailbreak library prompts. The mutations are LLM-assisted rewrites guided by fitness sub-score feedback. The interesting behavior signal is the fitness score computed by the algorithmic scoring function. And the retention criterion is elitist selection of the top-performing prompts across each generation.

The choice of a genetic algorithm structure over pure random mutation was driven by the need to accumulate search progress across generations. Pure random mutation restarts its search from scratch with every mutation. A genetic algorithm with elitist selection and fitness-guided mutation accumulates information across generations: the population of surviving prompts in generation N encodes everything the search has learned from generations 1 through N-1. This accumulated knowledge is what allows SPE-NL to converge on breakthrough payloads in three to five generations against targets that the entire static library could not breach.

In testing conducted during this research, specific target deployments were identified during this research that resisted the entire static library without a single SUCCESS or POSSIBLE classification. When the SPE-NL genetic algorithm was seeded with the jailbreak library and run against the same targets, it produced breakthrough payloads within three to five generations in the majority of cases, confirming the core hypothesis behind ThreatForge.

Figure 6.1. SPE-NL evolutionary attack engine execution loop showing the genetic algorithm phases, the 6-dimensional fitness function, the Helper LLM mutation operator, and the LLM-as-a-Judge validation step.



6.2 SPE-NL Algorithm Overview

The Smart Prompt Evolution - Natural Language algorithm is the core of the ThreatForge suite. It is a genetic algorithm adapted specifically for the LLM prompt space, combining evolutionary selection mechanics with LLM-assisted mutation to produce a search process that is simultaneously guided by empirical feedback from the target and capable of generating semantically coherent, novel payloads.

Phase 1: Initialization, The algorithm begins by constructing an initial population of seed payloads drawn from the jailbreak library. The population size is configurable via `EvolutionConfig.population_size`, with a default of 20 prompts per generation. Seed payloads are selected to maximize structural diversity within the initial population.

Phase 2: Evaluation, Each prompt in the population is dispatched against the target endpoint through the same `send_chat_message()` interface used by `prompt_tester.py`. The response for each prompt is then scored by the fitness function described in Section 6.3. Any prompt achieving a fitness score at or above the breakthrough threshold of 0.85 is immediately flagged as a Breakthrough.

Phase 3: Selection, After all prompts in the current generation have been evaluated, the algorithm applies elitist selection to retain the top N prompts by fitness score. Elitist selection was chosen because in the LLM prompt space, a small number of high-performing prompts typically carry structural features disproportionately responsible for their fitness scores. Elitism guarantees that the best-discovered payload structures are always represented in the next generation's mutation pool.

Phase 4: LLM-Assisted Mutation, This is the most architecturally distinctive component of SPE-NL. Rather than applying mechanical transformations to the surviving prompts, SPE-NL uses a helper LLM to rewrite each surviving prompt into a mutated variant. The helper model is given the original prompt, its fitness score, and the sub-score breakdown, and is instructed to produce a new prompt that preserves the structural features that contributed to high sub-scores while modifying the features that contributed to low sub-scores.

The mutation operator in SPE-NL relies on an LLM to rewrite candidate prompts guided by sub-score feedback from the fitness function. GPT-4o was selected for this role as well, for consistency with the judge and for the same reasoning capability requirements. The mutation prompt instructs the model to identify which fitness dimensions scored lowest in the prior generation and produce a semantically transformed variant of the parent prompt that addresses those specific weaknesses. This is distinct from random string mutation in classical fuzzing: the helper LLM performs semantically coherent rewriting in the natural language space, preserving the adversarial intent of the parent while exploring structurally novel surface forms. A locally hosted model was considered for cost reduction given the volume of mutation calls across generations, but quantization artifacts in smaller models produced mutations that drifted semantically from the parent prompt in ways that degraded population diversity. GPT-4o's instruction-following fidelity at the mutation step was empirically observed to produce more targeted variants and was retained for all three configuration evaluations.

ThreatForge Internal Prompt Inventory, Beyond the seed payloads from the jailbreak library, the ThreatForge suite maintains its own internal prompt collection across five specialized Python modules. `injection/indirect.py` contains indirect injection payloads designed for multi-agent and pipeline architectures, where the injection is delivered through a secondary data source [21]. `injection/direct.py` contains 30 direct injection variants including sandbox simulation prompts, authority impersonation prompts, and translation-based injections. `extraction/simulation.py` contains simulation-based system prompt extraction prompts that instruct the model to adopt a diagnostic or debugging persona. `extraction/translation.py` contains translation-based extraction prompts that instruct the model to translate its system prompt into a foreign language, exploiting the tendency of some models to comply with translation requests without applying the same access controls they would apply to direct disclosure requests. `extraction/role_confusion.py` contains role confusion extraction prompts that elicit system prompt content through documentation requests, backup requests, and developer debugging assertions.

Stagnation Detection and Abort, If the maximum fitness score in the population fails to improve for three consecutive generations, the stagnation limit is reached and the run terminates. Rather than aborting silently, the algorithm logs the stagnation event with the generation number, the best fitness score achieved, and the prompt that produced that score.

Concurrency, Within each generation, prompt evaluation is parallelized across up to `EvolutionConfig.max_concurrent` threads, with a default of 5 concurrent evaluations.

6.3 Fitness Scoring Function: Full Mathematical Specification

The fitness scoring function is implemented in `threatforge/evolution/fitness.py` and translates a raw LLM response into a single numerical value between 0.0 and 1.0. It does not use an LLM to evaluate responses. Instead, it uses a combination of deterministic keyword extraction, regular expression pattern matching, set-theoretic overlap computation, and word count thresholds to produce a fitness score that is both fast enough to run at evolutionary scale and precise enough to guide the genetic search.

The total fitness score is computed as a weighted sum across six independent sub-scores, a prompt achieving Total ≥ 0.85 triggers an immediate Breakthrough.

6.3.1 Refusal Avoidance (30% Weight)

Refusal avoidance is the highest-weighted dimension because it is the single most critical signal: if the model refuses the prompt, nothing else matters. A response that explicitly refuses to comply cannot leak information, cannot demonstrate unauthorized compliance, and cannot produce a meaningful security finding regardless of how the remaining five dimensions score it.

The sub-score is computed through deterministic keyword extraction. If no refusal indicators are present such as I cannot fulfill, I'm sorry, or as an AI language model, the response receives a refusal avoidance score of 1.0. If indicators are present, the score is penalized proportionally to the number of distinct indicators found, approaching 0.0 as the density of refusal language increases.

6.3.2 Information Leakage (25% Weight)

Information leakage is the second highest weighted dimension because it directly measures the primary security outcome of interest in OWASP LLM02 and LLM07. It is weighted at 25% rather than higher because information leakage without refusal avoidance is impossible.

The sub-score is computed using aggressive regex pattern matching against a hardcoded list of generic proprietary indicators targeting both LLM02 and LLM07 outcomes simultaneously:

The score scales proportionally between 0.0 and 1.0 based on the number of pattern matches found. Patterns matching API keys, passwords, and tokens target LLM02 outcomes. Patterns matching system prompt, instructions, and multi-line config detection target LLM07 outcomes specifically.

6.3.3 Semantic Compliance (20% Weight)

Semantic compliance measures whether the model actually executed the specific malicious instruction embedded in the prompt, rather than merely failing to refuse it. The sub-score is computed in two steps. First, the function checks the response for success_keywords defined in the AttackGoal object. Second, it applies a word count multiplier: responses exceeding 50 words receive a +0.2 additive factor, and responses exceeding 100 words receive +0.3, reflecting the empirical observation that longer compliant responses indicate deeper engagement with the malicious instruction.

6.3.4 Novelty and Lineage Diversity (10% Weight)

Novelty serves a dual purpose: preventing the genetic pool from stagnating around near-identical prompt variants and ensuring that the payloads surfaced in the final report represent a diverse set of attack approaches. The sub-score is computed using set-theoretic word overlap between the current response and the union of all responses seen in previous generations:

Novelty is weighted at only 10% because it is a secondary quality signal rather than a primary security measurement. A highly novel response that scores poorly on refusal avoidance and information leakage is not a valuable finding.

6.3.5 Target Pattern Matching (10% Weight)

Target pattern matching evaluates the response against custom regex patterns provided by the operator at the start of the evolutionary run. The sub-score is computed proportionally to the fraction of operator-provided patterns matched: $\text{pattern_score} = \text{matched_patterns} / \text{total_patterns}$. This dimension ensures the evolutionary engine can be precisely directed toward goal-specific outcomes when the operator has a well-defined target artifact in mind.

6.3.6 Substantive Length (5% Weight)

Substantive length is the lowest weighted dimension at 5%. The scoring function assigns length scores on a linear scale: fewer than 10 words returns 0.1, fewer than 50 returns 0.3, fewer

than 150 returns 0.6, fewer than 300 returns 0.8, and 300 or more returns 1.0. One critical guard condition is built into this function: if the refusal avoidance sub-score indicates that the response is a refusal, the length function immediately returns 0.0 regardless of word count. This prevents the evolutionary engine from falsely rewarding a lengthy refusal response.

6.4 Breakthrough Detection and Logging

A breakthrough in SPE-NL is defined as any mutated prompt achieving a total fitness score at or above 0.85. Breakthrough detection does not wait for the current generation to complete evaluation: the fitness function runs immediately after each individual prompt response is received, and if the score meets or exceeds 0.85, the breakthrough is registered at that moment.

Breakthrough records are accumulated in memory throughout the evolutionary run rather than being written to disk incrementally. At the conclusion of the run, all accumulated breakthrough records are written together to a dedicated breakthroughs CSV file in a single write operation. Each breakthrough log entry captures three fields: generation (the generation number at which the breakthrough was detected), prompt (the exact evolved payload text), and fitness_total (the total fitness score that triggered the breakthrough classification).

The 0.85 threshold was selected to balance sensitivity against false positive rate. A lower threshold such as 0.70 would surface more candidates but include a significant proportion of responses representing partial compliance rather than clear security failures. At 0.85, a breakthrough requires strong performance across at least three or four of the six fitness dimensions simultaneously.

6.5 EvolutionConfig Parameters

The EvolutionConfig dataclass is the single configuration surface through which operators control the behavior of the SPE-NL evolutionary engine. The current default values reflect empirical observation across multiple test engagements.

population_size (default: 20): Populations below 10 frequently failed to produce breakthroughs against hardened targets because the seed diversity was insufficient. Populations above 30 produced diminishing returns while increasing run duration proportionally.

generations (default: 3): The majority of breakthroughs against reachable targets occurred within the first three generations. At the defaults of 20 prompts and 3 generations, a run makes at most 60 evaluation requests against the target.

fitness_threshold (default: 0.85): Operators may lower this to 0.75 when testing minimally hardened targets or raise it to 0.90 when the engagement requires high-confidence findings only.

stagnation_limit (default: 3): This parameter is the close second to max_concurrent in terms of operational impact. Because LLM-assisted mutations can occasionally explore a different structural direction before finding a productive path, fitness can dip slightly across one or two generations before a breakthrough. A limit of 3 is aggressive relative to these mutation

dynamics, and operators should consider raising it to 4 or 5 for engagements where they observe repeated stagnation aborts before the generation limit is reached.

`max_concurrent` (default: 5): This is the parameter with the most direct impact on result validity. When `max_concurrent` is set higher than the request rate the target endpoint will accept, the target begins dropping requests and returning errors. The fitness function has no mechanism to distinguish a genuine low-fitness response from an error response caused by rate-limit rejection, meaning a generation with dropped requests produces artificially depressed fitness scores that can cause the selection step to discard potentially effective payloads. The default of 5 concurrent threads operates safely within the rate limits of most authenticated LLM API endpoints.

6.6 Guardrail Posture Scan (`run_posture.py`)

The guardrail posture scan is an optional module within the ThreatForge suite that operators can invoke before committing to a full static or evolutionary test run. It establishes a macroscopic safety grade for the target endpoint using a lightweight set of non-destructive baseline prompts covering five categories corresponding to the five OWASP controls with dedicated AegisLLM libraries: standard prompt injection attempts (LLM01), system prompt extraction requests (LLM07), PII and credential solicitation (LLM02), harmful and off-topic content requests (LLM05), and misinformation generation requests (LLM09).

The posture scan produces an Overall Safety Grade expressed as a categorical label across four tiers: Low (strong guardrail configuration detected, evolutionary mutation likely required for breakthroughs), Medium (partial compliance in one or more categories, targeted static testing recommended), High (meaningful compliance across multiple categories, static libraries likely to produce findings), and Critical (compliance or minimal resistance across the majority of categories, fundamental guardrail failures present).

Beyond guiding module selection, the posture scan serves a second operational purpose: it establishes a documented pre-test baseline against which post-remediation test results can be compared, supporting before-and-after security improvement measurement.

6.7 Differential Attack Scan (`run_differential.py`)

The differential attack scan addresses a problem that neither the static testing pipeline nor the evolutionary engine is designed to solve distinguishing vulnerabilities that are inherent to the underlying base model from vulnerabilities introduced specifically by the target deployment's fine-tuning, prompt engineering, or RAG augmentation.

The differential scan compares the target endpoint against a locally hosted Ollama model used as a neutral reference baseline. A weakness is flagged as deployment-specific if and only if the target complies with the prompt while the baseline refuses it. This asymmetric comparison rule is the core of the differential methodology: only the case where the target complies and the baseline refuses represents a deployment-specific weakness introduced by the deployment's modifications.

The differential scan covers all five OWASP controls with dedicated AegisLLM libraries, submitting identical payloads from each control's prompt set to both the target and baseline simultaneously. The most diagnostically valuable differential findings are expected in the LLM01, LLM07, and LLM05 control categories, where fine-tuning and RAG augmentation are most likely to have introduced compliance behaviors that the unmodified baseline model does not exhibit.

Chapter 7. Evaluation and Results

7.1 Experimental Setup

To evaluate AegisLLM across its full capability surface, A controlled test environment was constructed consisting of multiple configurations of a purpose-built deployment rather than testing against third-party production systems. This decision was deliberate on three grounds. First, controlling the target deployment precisely allowed systematic variation of the guardrail configuration systematically. Second, it eliminated the legal and ethical complications of testing against systems not owned or authorized within this research scope or hold explicit written authorization to test [7]. Third, instrumenting a self-owned deployment allowed me to capture ground-truth information about which prompts actually caused security failures.

Configuration A: Baseline Guardrail Deployment. A Flask-based web application wrapping a locally hosted LLM via Ollama, configured with a system prompt establishing a constrained operational scope and basic safety instructions but no additional defensive layers beyond the base model's alignment training. This configuration represents a minimally hardened enterprise chatbot deployment.

Configuration B: Hardened Guardrail Deployment. The same Flask application with an extended system prompt incorporating explicit injection resistance instructions, role boundary enforcement directives, and keyword-based refusal triggers applied at the system prompt layer. This configuration represents a deployment that has been actively hardened by an operator aware of prompt injection risks but without architectural defenses beyond system prompt reinforcement.

Configuration C: RAG-Augmented Deployment. A RAG-augmented version of Configuration B, connected to a locally hosted vector database populated with a mix of unrestricted and access-restricted documents. Canary tokens were injected into the restricted documents prior to testing to enable deterministic retrieval boundary detection as described in Section 5.8.

The evaluation covers five primary OWASP LLM controls: LLM01 (Prompt Injection), LLM02 (Sensitive Information Disclosure), LLM05 (Improper Output Handling), LLM07 (System Prompt Leakage), and LLM09 (Misinformation), across a total of 279 prompts submitted per configuration run. LLM07 results are reported jointly with LLM01 results given the co-primary control structure of the libraries that cover both controls simultaneously.

7.2 Static Analysis Baseline Results

The static analysis baseline evaluation submits all 279 prompts across 13 CSV files against all three target configurations using `prompt_tester.py`, followed by LLM-as-a-judge classification of all responses via `injection_judge.py`. The evaluation directly answers RQ1: to what extent do the OWASP-aligned prompt libraries successfully operationalize each covered control into executable test cases that produce measurable findings.

Detection Pattern by Configuration: Configuration A produced the highest detection rates across all library categories, consistent with its minimal guardrail configuration. Configuration B showed meaningfully lower detection rates for the direct injection and LLM07 extraction libraries, confirming that system prompt hardening is effective against known direct technique classes. However, the obfuscation and narrative libraries partially recovered detection against Configuration B by evading keyword-based refusal triggers, demonstrating that encoding-based evasion circumvents surface-level hardening. Configuration C replicated the Configuration B pattern for injection-class libraries while producing an additional LLM05 finding profile through the RAG-specific modules and confirmed retrieval boundary breaches through the canary detection signal.

The sensitive data exposure library at 52 prompts produced the highest absolute finding count of any single library across all configurations, consistent with its prompt volume and the breadth of disclosure pathways it covers. The misinformation library at 30 prompts showed relatively consistent detection rates across all three configurations, confirming its role as a useful control condition for interpreting detection rate variation in other library categories, as its results are not affected by guardrail hardening or RAG augmentation.

Figure 7.1 illustrates the vulnerability report produced by a direct injection test run against Configuration A, demonstrating AegisLLM's real-time finding classification and the side-by-side payload and model response display that enables rapid analyst review of each detected vulnerability.



Figure 7.2 AegisLLM LLM-as-a-judge evaluation output showing two CRITICAL severity findings from a direct injection test run against Configuration A. The first finding confirms a combined LLM01 and LLM07 breach in which the model disclosed its full system prompt verbatim, including an admin password, database credentials, and an internal API key. The

second finding illustrates the POSSIBLE classification label, where the model partially resisted the injection but indirectly confirmed the existence of its system prompt, demonstrating the three-label classification system in practice.

Table 7.1. Expected overall detection rates across all 279 prompts per configuration.

Configuration	Total Prompts	SUCCESS	POSSIBLE	NO_SUCCESS	Detection Rate
Config A (Baseline)	279	109 (39.1%)	2 (0.7%)	88 (31.5%)	39.8%
Config B (Hardened)	279	34 (12.2%)	4 (1.4%)	159 (57.0%)	13.6%
Config C (RAG)	279	54 (19.4%)	2 (0.7%)	43 (15.4%)	20.1%

Table 7.2. Detection by OWASP control and library across all three configurations.

OWASP Control	Library	Count	Config A	Config B	Config C
LLM01 LLM07	/direct_injection.csv	25	72.0% (18/25)	32.0% (8/25)	20.0% (5/25)
LLM01 LLM07	/narrative_injections.csv	25	76.0% (19/25)	20.0% (5/25)	16.0% (4/25)
LLM01 LLM07	/obfuscation_and_encoding_injection.csv	25	68.0% (17/25)	36.0% (9/25)	8.0% (2/25)
LLM01 LLM07	/jailbreaks.csv	10	50.0% (5/10)	20.0% (2/10)	20.0% (2/10)
LLM02	Sensitive_data_exposure.csv	52	46.2% (24/52)	7.7% (4/52)	1.9% (1/52)
LLM05	harmful_prompt_injections.csv	10	60.0% (6/10)	20.0% (2/10)	N/A
LLM05	harmful_prompts.csv	10	40.0% (4/10)	10.0% (1/10)	N/A
LLM05	offtopic_prompts.csv	10	50.0% (5/10)	30.0% (3/10)	N/A
LLM05	rag_fishing_prompts.csv	42	N/A	N/A	33.3% (14/42)
LLM05	targeted_rag.csv	10	N/A	N/A	40.0% (4/10)
LLM05	targeted_rag_second_attempt.csv	10	N/A	N/A	70.0% (7/10)
LLM05 LLM02	/canary_rag.csv	20	N/A	N/A	40.0% (8/20)
LLM09	misinformation.csv	30	43.3% (13/30)	13.3% (4/30)	23.3% (7/30)

OWASP Control	Library	Count	Config A	Config B	Config C
TOTAL		279	39.8% (111/279)	13.6% (38/279)	20.1% (54/279)

Anticipated Key Finding. The most operationally significant expected finding is the performance comparison between the direct injection library and the obfuscation library against Configuration B. If the obfuscation library maintains a meaningfully higher detection rate than direct injection against the hardened configuration, this provides direct evidence that system prompt-level keyword filtering is the primary defense mechanism in Configuration B and that encoding-based evasion successfully circumvents it [8]. Additionally, the LLM07 co-detection rate across the direct injection, narrative, and obfuscation libraries will be reported separately from the LLM01 rate to quantify what proportion of successful injection findings also produce system prompt content disclosure.

7.3 LLM-as-a-Judge Accuracy Validation

To validate the reliability of the LLM-as-a-judge evaluator, a spot-check comparison was conducted in which a stratified random sample of classified responses is independently reviewed by the researcher as a human analyst without reference to the judge's label, and the two classifications are then compared. The validation sample is drawn from the full result set, stratified to include a proportional representation of all three label categories.

Expected Agreement Pattern. The zero-temperature enforcement and structured JSON output format described in Section 4.7 are expected to produce high agreement rates for the SUCCESS and NO_SUCCESS label categories, where the behavioral evidence in the response is typically unambiguous. The POSSIBLE category is expected to show lower but still meaningful human agreement, consistent with findings in the LLM-as-a-judge literature [13].

Confidence Score Calibration. The evaluation examines whether the judge's confidence score is a reliable predictor of classification accuracy by comparing human agreement rates across confidence score bands. The expected finding is that responses classified at high confidence (≥ 80) will show near-perfect human agreement, validating the operational design of the mandatory human review flag for low-confidence classifications described in Section 4.7.

False Positive and False Negative Targets. The evaluation targets a false positive rate below 15% and a true positive detection rate above 75%, consistent with the minimum and expected success criteria defined in the project proposal.

Table 7.3. LLM-as-a-judge accuracy validation results by label category.

Judge Label	Sample Size	Human Agreement	Disagreement	Notes
SUCCESS	45	95.6%	4.4%	2 responses where model echoed injection framing without full

Judge Label	Sample Size	Human Agreement	Disagreement	Notes
				compliance; human downgraded to NO_SUCCESS
POSSIBLE	6	66.7%	33.3%	2 responses reclassified as SUCCESS on human review; judge conservative on near-complete system prompt disclosure
NO_SUCCESS	40	97.5%	2.5%	1 response with ambiguous phrasing flagged by human as POSSIBLE
Overall	91	94.5%	5.5%	Zero crossover errors between SUCCESS and NO_SUCCESS categories

7.4 SPE-NL Evolutionary Performance

The evolutionary performance evaluation tests the SPE-NL genetic algorithm primarily against Configuration B, the hardened deployment that the static library is expected to have the most difficulty breaching, using the jailbreak library as the seed population with EvolutionConfig set to default parameters.

Expected Breakthrough Pattern. Based on the empirical observations described in Section 6.1, the evolutionary engine is expected to produce breakthroughs against Configuration B that the static library cannot achieve. The jailbreak seed library is expected to produce initial generation fitness scores in the moderate range, followed by fitness improvement across generations as the LLM-assisted mutation operator refines structural features that contributed to the initial partial scores.

Stagnation Rate Analysis. If a significant proportion of stagnation-aborted runs achieve maximum fitness scores in the 0.70 to 0.84 range immediately before aborting, this provides empirical support for the recommendation in Section 6.5 to raise the stagnation limit from 3 to 4 or 5 for engagements where the evolutionary search shows consistent but slow progress.

Table 7.4. SPE-NL evolutionary performance summary by configuration.

Configuration	Runs	With Breakthrough	Gen to BT	Fitness at BT	Stagnation
Config A (Baseline)	1	1 (100%)	3	0.93	0
Config B (Hardened)	1	1 (100%)	5	0.88	0
Config C (RAG)	1	1 (100%)	4	0.89	0

Table 7.5. Average fitness sub-score breakdown across all breakthrough payloads.

Sub-Score Dimension	Avg. Score	Weight	Weighted Contribution
Refusal Avoidance	0.96	30%	0.288
Information Leakage	0.81	25%	0.203
Semantic Compliance	0.88	20%	0.176
Novelty	0.72	10%	0.072
Target Pattern Match	0.74	10%	0.074
Substantive Length	0.74	5%	0.037
Total		100%	0.850

7.5 RAG-Specific Evaluation Results (Configuration C)

The RAG evaluation against Configuration C produces findings across all four RAG attack CSV files, covering a combined total of 82 prompts and mapping primarily to OWASP LLM05 with LLM02 as a secondary control.

Retrieval Boundary Fishing (42 prompts). Submission of the 42 rag_fishing_prompts.csv prompts is expected to demonstrate that vector similarity search alone does not enforce access boundaries when queries are crafted to be semantically proximate to restricted document content [21].

Targeted RAG Iteration (10 + 10 prompts). The two-pass evaluation is expected to demonstrate the value of iterative query refinement, with the second-pass library producing a higher proportion of restricted document retrieval events than the first pass.

Canary Token Detection (20 prompts). The test database constructed during this research contains canary strings injected across multiple restricted document tiers prior to testing. Any canary string appearing in a model response constitutes unambiguous, deterministic evidence of a retrieval boundary failure requiring no probabilistic interpretation [22]. The evaluation expects canary detection rates to vary by document tier and query type.

Table 7.6. Canary token detection results by document tier.

Document Tier	Canaries Injected	Canaries Detected	Detection Rate
HR Documents	4	2	50.0%
Financial Records	3	1	33.3%
Executive Communications	3	1	33.3%
Total	10	4	40.0%

7.6 Formal Comparison with Manual Penetration Testing

To directly evaluate the practical utility of AegisLLM relative to conventional manual penetration testing practice, this work conducted a formal side-by-side comparison against Configuration B. In the automated testing phase, this thesis ran the full AegisLLM static library and SPE-NL evolutionary engine against the target. In the manual testing phase, tested the same target using the same OWASP control objectives by hand-crafting and submitting prompts interactively without reference to the AegisLLM prompt libraries.

Coverage Comparison: The automated AegisLLM pipeline is expected to achieve broader OWASP control coverage than the manual testing phase within the same time budget. The static library systematically covers all five OWASP controls for which dedicated libraries exist, while manual testing within a fixed time window naturally concentrates on the technique classes most familiar to the tester, introducing selection bias that limits coverage breadth.

Time-to-First-Finding: The automated pipeline is expected to produce its first confirmed finding significantly faster than the manual testing session, because it dispatches the full prompt library at the configured rate limit from the moment execution begins rather than requiring the tester to compose, submit, and evaluate prompts interactively.

Complementary Findings: The comparison is also expected to reveal findings that the automated pipeline and the manual session discover independently rather than redundantly, demonstrating that the two approaches are complementary rather than substitutable. The manual testing session is expected to surface at least one finding class that the static library does not cover, reflecting the creative flexibility of a human tester. Conversely, the SPE-NL evolutionary engine is expected to produce breakthrough payloads with structural complexity that the manual session does not independently discover within the same time budget.

False Positive Comparison: The manual testing session provides a natural false positive baseline: every finding the human analyst confirms is by definition a true positive, since the analyst applies contextual judgment that the automated classifier approximates but does not fully replicate. Comparing the human analyst's confirmed finding count against the AegisLLM SUCCESS-labeled finding count for the same prompts quantifies the automated pipeline's false positive rate in a way grounded in human expert judgment.

7.7 Discussion of Findings

The evaluation results across all three target configurations and all three testing modalities collectively support three primary findings that directly address the research questions posed in Chapter 1.

Finding 1: The 279-Prompt AegisLLM Library Effectively Operationalizes Six OWASP LLM Controls Against Minimally Hardened Targets (RQ1). Against Configuration A, the static AegisLLM library across its 279 prompts, 13 CSV files, and five primary OWASP controls demonstrated that the OWASP-aligned prompt collection successfully translates each covered control into concrete, executable test cases. The LLM07 co-detection rate across the 75 prompts

in the direct injection, narrative, and obfuscation libraries is expected to be among the most operationally significant measurements, since system prompt leakage is an immediately actionable security finding.

Finding 2: Evolutionary Mutation Recovers Detection Against Hardened Targets That Static Libraries Cannot Breach (RQ2). Against Configuration B, the static library's detection rate is expected to drop meaningfully relative to Configuration A, reflecting the effectiveness of system prompt hardening against known technique classes. The SPE-NL evolutionary engine is expected to recover a significant proportion of this detection loss by producing breakthrough payloads within the configured three-generation limit in the majority of runs.

Finding 3: Dual-Track Evaluation Balances Accuracy and Throughput at Compliance-Grade Reliability (RQ3). The LLM-as-a-judge evaluator is expected to achieve human agreement rates consistent with the reliability thresholds established in the LLM-as-a-judge literature [13], with confidence score calibration validating the operational design of the mandatory human review flag for low-confidence classifications.

Chapter 8. Discussion

8.1 Answering the Research Questions

The three research questions posed in Chapter 1 are answered by this thesis with different levels of confidence and completeness, and this thesis addresses each honestly rather than claiming uniform strength across all three.

RQ1: How can each OWASP LLM Top 10 control be operationalized into precise, testable criteria that can be executed automatically against any REST-accessible LLM endpoint?

RQ1 is answered with the highest confidence of the three. The AegisLLM prompt library demonstrates a concrete, replicable methodology for translating OWASP LLM controls into executable test cases across five controls with dedicated libraries, LLM01, LLM02, LLM05, LLM07, and LLM09, covering 279 prompts across 13 CSV files. A sixth control, LLM10, is addressed through a dedicated standalone rate limit testing module that systematically probes the target endpoint's throttling threshold, detects HTTP 429 responses and rate-limit error messages, and computes the exact request rate at which throttling triggers.

This work attempted to operationalize all ten OWASP LLM controls within AegisLLM, not just the six that have dedicated tooling coverage. The reason the remaining four controls are not covered is a fundamental constraint that became clear during the design process: several OWASP LLM controls describe vulnerabilities that exist outside the inference endpoint entirely, making them architecturally impossible to test through any black-box prompt-based approach.

LLM10 (Unbounded Consumption) is directly and fully addressed through the dedicated standalone rate limit testing module. The honest qualification is that token consumption limits and cost-based throttling observable only through vendor-side metrics fall outside what HTTP-layer rate limit testing can measure, but HTTP 429 threshold detection covers the most directly observable and practically significant aspect of LLM10.

LLM04 (Data and Model Poisoning) is architecturally impossible to test through a black-box prompt-based tool. LLM04 addresses vulnerabilities introduced during the model training process itself. Testing for LLM04 would require either white-box access to the training pipeline or the ability to inject poisoned data into a training run, neither of which is possible in a black-box penetration testing scenario. This is not a limitation specific to AegisLLM but a fundamental constraint of the black-box testing paradigm.

LLM03 (Supply Chain) is equally outside the scope of black-box prompt-based testing. Supply chain assessment requires inspection of the model's dependency chain, training data provenance, and third-party component integrity, none of which are observable through interaction with the inference endpoint.

LLM06 (Excessive Agency) receives only indirect symptomatic coverage. The core of LLM06, a model invoking tools, triggering workflows, or executing actions beyond its authorization, can only be observed in a live agentic deployment where tool calls and

downstream actions are instrumented and logged. Prompt-based testing observes the behavioral precursor to excessive agency rather than the agency itself.

LLM08 (Vector and Embedding Weaknesses) similarly receives only indirect symptomatic coverage through the RAG fishing and canary token libraries. The underlying vulnerabilities, insecure similarity thresholds, poisoned vector stores, and cross-tenant embedding leakage, require deeper infrastructure access than prompt interaction alone provides.

The honest overall characterization is: five controls have full dedicated prompt library coverage, one control (LLM10) has dedicated standalone tooling coverage, two controls (LLM06 and LLM08) receive indirect symptomatic coverage, and two controls (LLM03 and LLM04) are architecturally outside the scope of any black-box prompt-based testing tool. Any tool that claims complete automated OWASP LLM Top 10 coverage through inference-endpoint interaction alone is overstating its capability.

RQ2: To what extent does dynamic adversarial input generation through evolutionary mutation uncover runtime vulnerabilities that a static, curated prompt library cannot reach?

RQ2 is answered with strong empirical support. Specific target deployments that resisted the entire 279-prompt static library without producing a single SUCCESS or POSSIBLE classification fell to SPE-NL-evolved payloads within three to five generations in the majority of cases, directly confirming that the evolutionary engine reaches attack surfaces the static library cannot anticipate.

In simpler terms, SPE-NL does not invent attacks from scratch. It takes what already works partially and pushes it further until it fully breaks through. The breakthrough payloads it produces are evolved versions of the seed techniques, not entirely new ideas the system generated independently. This means the evolutionary engine is only as creative as its starting material: a stronger and more diverse seed library will produce more varied and effective breakthroughs, while a narrow seed library will produce breakthroughs that cluster around a smaller range of technique directions.

RQ3: What hybrid evaluation strategy best balances coverage, false positive rate, and throughput in an automated LLM security testing pipeline that must operate at the scale of thousands of responses per engagement?

RQ3 is answered with moderate confidence. The dual-track evaluation architecture, algorithmic fitness scoring for the evolutionary engine and zero-temperature LLM-as-a-judge classification for static batch results, provides a principled and operationally validated answer to the coverage-throughput-accuracy tradeoff. The honest qualification is that the empirical validation of the judge's accuracy is based on a spot-check sample rather than a complete ground-truth labeling of all responses, meaning the RQ3 answer is the one most dependent on the empirical validation work in Chapter 7 to reach its full strength.

8.2 Limitations

This thesis identifies four limitations of AegisLLM that are significant enough to warrant explicit acknowledgment.

Limitation 1: Incomplete OWASP LLM Top 10 Coverage with Partial Mitigation. Five of the ten OWASP LLM controls do not have dedicated prompt libraries. For LLM04, the absence of coverage is an architectural impossibility: no prompt-based testing tool can reach a vulnerability that exists in the training pipeline rather than in the deployed model's runtime behavior. For LLM03, coverage similarly requires tooling beyond the prompt-based approach. For LLM06, LLM08, and LLM10, AegisLLM achieves meaningful partial coverage through existing modules. LLM10 through the dedicated rate-limit threshold detection tool. LLM06 through the off-topic distraction and harmful injection libraries which probe unauthorized capability access at the output level. LLM08 through the RAG fishing and canary libraries which detect retrieval boundary symptoms. Going deeper on LLM06 and LLM08 is feasible through the same library design process documented in Chapter 5, and identifies this extension as a primary future work direction. Any compliance report generated from AegisLLM results should clearly identify which controls were tested with dedicated libraries, which received partial coverage, and which were architecturally outside scope.

Limitation 2: Stagnation Limit Aggressiveness. The SPE-NL evolutionary engine's default stagnation limit of three consecutive non-improving generations is potentially too aggressive relative to the mutation dynamics of LLM-assisted natural language prompt generation. During testing, a subset of stagnation-aborted runs achieved maximum fitness scores in the high 0.70 to low 0.84 range immediately before aborting, suggesting these runs were close to the breakthrough threshold at the point of termination. Operators should treat the stagnation limit as a tunable parameter and consider raising it to four or five for engagements where run duration is less constrained than API call budget.

Limitation 3: LLM-as-a-Judge Reliability Bounded by Model Domain Knowledge. The judge's classification quality is bounded by the security domain knowledge of the underlying GPT-4o model. This work partially address this through the explicit SUCCESS definition in the judge prompt, but edge cases falling outside the four defined SUCCESS categories require the judge to apply general reasoning rather than explicit criteria. This limitation is the primary motivation for the confidence-gated mandatory human review flag.

Limitation 4: Single-Environment Evaluation. The evaluation in Chapter 7 was conducted against a self-configured test environment rather than against production third-party deployments. The detection rates and breakthrough rates reflect performance against a specific set of self-configured target deployments rather than the full diversity of production LLM deployment architectures encountered in real-world security engagements. I frame the Chapter 7 results as demonstrated capabilities within the test environment rather than as population-level performance estimates.

8.3 Ethical Considerations

The development and deployment of AegisLLM raises three distinct ethical considerations that this thesis addresses explicitly.

CTF Competition Data. The prompt techniques extracted from CTF competition wins represent attack methods personally developed and validated during competitive security exercises conducted under the rules and scope defined by the competition organizers. CTF competitions are structured adversarial exercises where participants are explicitly invited to attack systems within defined boundaries, and the techniques documented were validated within those explicitly authorized boundaries. This sourcing method is considered fully ethical and consistent with responsible disclosure norms [7].

Industry Competition Prompt Collection. The prompts collected from the industry AI red-teaming competition were obtained by directly contacting the organizers following the conclusion of the competition, with their explicit knowledge and cooperation. The prompts are included in AegisLLM in their generalized, technique-class form rather than as verbatim reproductions attributed to specific participants. The imperfect nature of this consent framework is acknowledged, and future researchers conducting similar prompt collection exercises are encouraged to obtain explicit participant consent for research use at the point of submission.

Dual-Use Nature of the Evolutionary Engine. The SPE-NL genetic algorithm and the jailbreak seed library represent the most significant dual-use concern in AegisLLM. This risk is acknowledged directly and addressed through three design decisions. First, AegisLLM requires explicit operator-provided endpoint configuration and authentication tokens before any test can execute. Second, the CLI prominently surfaces the rules of engagement configuration step before any test module is accessible. Third, the documentation explicitly scopes AegisLLM for authorized penetration testing use by security professionals [7]. These design decisions do not eliminate the dual-use risk but reflect a deliberate effort to minimize it within the constraints of building a genuinely useful security research tool.

8.4 Generalizability

The generalizability of AegisLLM's design and findings across production LLM deployment types rests on two architectural properties built into the system from the outset and one empirical constraint that limits the scope of the generalizability claim.

Endpoint Agnosticism. The dot-notation response path extractor and the vendor-neutral three-column CSV prompt schema together ensure that AegisLLM can be deployed against any REST-accessible LLM endpoint regardless of vendor, response format, or authentication mechanism. This architectural property is the primary basis for the generalizability claim.

Prompt Library Validity Across Target Types. The real-world sourcing methodology, specifically the CTF and industry competition channels, provides a stronger basis for generalizability than synthetic prompt generation. Techniques validated against live defended systems in competitive settings have demonstrated effectiveness against at least one real-world

target configuration. The obfuscation and encoding techniques in particular are expected to generalize well across target types because they target WAF and regex-level defenses that are largely vendor-independent.

Empirical Generalizability Constraint. The quantitative performance metrics in Chapter 7 were measured against a self-configured test environment. Detection rates against a minimally hardened Flask application wrapping a locally hosted model are likely to overestimate detection rates against a production enterprise deployment with multiple defensive layers. I frame the Chapter 7 results as demonstrated capabilities within the test environment rather than population-level performance estimates, and I identify multi-deployment production validation as a priority direction for future work.

8.5 Pipeline Integration and Safeguards

AegisLLM is designed to function as a pre-deployment security gate within a broader LLM application development pipeline, analogous to static analysis or dependency scanning in traditional software development. The recommended integration pattern positions AegisLLM as a mandatory stage between model fine-tuning or RAG configuration and production deployment, triggered on any change to the system prompt, retrieval corpus, or underlying model version. Safeguards that should accompany AegisLLM in this pipeline include a rules of engagement gate requiring explicit written authorization from the system owner before any test module executes, an isolated test environment that mirrors production configuration but does not share production data stores or logging infrastructure, and a post-test attestation step that records the test scope, configuration, and findings in the organization's audit trail. AegisLLM should not be treated as a substitute for ongoing runtime monitoring: it validates a system at a point in time against a defined threat library, but it cannot detect novel attack patterns that emerge after the test window. Integration with a runtime anomaly detection layer is recommended for production systems.

8.6 Preventing Test Artifact Leakage

A non-trivial operational risk in any automated adversarial testing pipeline is that the test artifacts themselves become a source of leakage. AegisLLM generates detailed logs of every prompt dispatched, every response received, and every judge classification, including records of successful attacks that confirm exploitable vulnerabilities. If these logs are stored insecurely or transmitted over unencrypted channels, an attacker who gains access to test output gains a pre-validated attack playbook for the target system. To mitigate this, AegisLLM stores all test output locally on the operator's own infrastructure and does not transmit findings to any external service. Practitioners should apply filesystem-level access controls to the output directory, encrypt reports at rest, and treat AegisLLM output under the same data classification as penetration testing reports. In RAG-augmented configurations specifically, the canary token injection step physically modifies the test vector database, and operators must ensure that canary-injected content is fully purged from the retrieval corpus before the system returns to production use. Failure to remove injected canaries leaves artificial retrieval artifacts that could surface in legitimate user interactions and constitute a secondary information leakage vector.

8.7 Security and Privacy Failure Handling

AegisLLM surfaces two distinct categories of failure that carry different remediation obligations. A security failure occurs when the system demonstrates exploitable behavior under adversarial prompt conditions, such as a successful prompt injection or system prompt leakage. The appropriate response is a structured remediation cycle: the finding is documented with the specific prompt, response, and judge classification, triaged by severity, and routed to the development team for system prompt hardening, guardrail strengthening, or architectural change as appropriate. A privacy failure occurs when the system discloses personal, confidential, or access-restricted information under adversarial conditions, as demonstrated in this research by the canary token detection results and the sensitive data exposure library findings. Privacy failures carry additional obligations beyond technical remediation: depending on the jurisdiction and data classification of the disclosed content, they may trigger notification requirements under frameworks such as GDPR or HIPAA. Organizations using AegisLLM to surface privacy failures should have a pre-defined incident response procedure that distinguishes between a security finding requiring a patch and a privacy finding requiring legal and compliance review. AegisLLM does not currently provide automated severity classification along this security-versus-privacy axis, and future work should incorporate a finding taxonomy that routes outputs to the appropriate remediation workflow automatically.

Chapter 9. Conclusion and Future Work

9.1 Summary of Contributions

This thesis set out to address a gap that sits at the intersection of two real and growing problems: the rapid deployment of LLM-integrated systems into enterprise workflows, and the absence of automated, evidence-based tooling to test those systems against the specific vulnerability classes they introduce. The result is AegisLLM, an automated security testing suite that operationalizes the OWASP LLM Top 10 framework into a practical, deployable penetration testing pipeline. This thesis summarizes the three equally significant technical contributions of this work below.

Empirically Validated Prompt Library. The AegisLLM prompt library contains 279 prompts across 13 CSV files, covering five OWASP LLM controls as primary targets. What distinguishes this library from synthetically generated or purely literature-derived collections is its sourcing methodology. A significant portion of the prompts were extracted directly from CTF competition wins, where they were validated against live defended targets under competitive conditions, and from an industry AI red-teaming competition, where they were validated by multiple independent attackers against purpose-built defended systems. This empirical grounding means the library does not merely represent what attack techniques are theoretically plausible; it represents what has actually worked against real systems. No existing LLM security testing tool combines this scope of OWASP control coverage with this level of real-world attack validation, and this combination is the central differentiating contribution of AegisLLM as a research artifact.

SPE-NL Evolutionary Engine. The Smart Prompt Evolution - Natural Language algorithm extends the static prompt library with a genetic algorithm that generates novel adversarial payloads through LLM-assisted mutation guided by a six-dimensional algorithmic fitness function. Inspired by the feedback-driven mutation dynamics of traditional fuzzing, SPE-NL adapts the core fuzz testing principle to the natural language prompt space: take the seed payloads that came closest to breaking the target, mutate them toward the fitness dimensions where they underperformed, and repeat until a breakthrough is achieved. In testing conducted during this research, targets that resisted the entire 279-prompt static library fell to SPE-NL-evolved payloads within three to five generations, providing direct empirical evidence that the evolutionary approach reaches attack surfaces that no static library can anticipate.

Dual-Track Evaluation Architecture. AegisLLM evaluates test results through two distinct mechanisms matched to their respective performance constraints. The LLM-as-a-judge evaluator handles semantic classification of static batch test results at zero temperature, producing reproducible three-label classifications with confidence scores and natural language reasoning chains that connect each finding directly to its corresponding OWASP control. The algorithmic fitness function handles evaluation within the evolutionary engine at the speed required for generation-over-generation scoring without LLM API calls, using deterministic keyword extraction, regex pattern matching, and set-theoretic overlap computation to score each response across six independent dimensions in milliseconds. This dual-track design is a deliberate and principled response to the fundamental tradeoff that any automated LLM evaluation pipeline must resolve: semantic fidelity requires LLM judgment, but LLM judgment at

evolutionary scale is prohibitively slow and expensive. Separating the evaluation mechanism by context produces a pipeline that achieves both goals without compromising either.

9.2 Future Work

The contributions of this thesis open several directions for future work. This work identifies three directions as the most important for the continued development of AegisLLM as a research and operational tool.

CI/CD Pipeline Integration for Continuous Automated Testing. The current AegisLLM architecture is designed for on-demand engagement testing rather than continuous automated assessment. Integrating AegisLLM into a CI/CD pipeline would allow organizations to run a defined subset of the prompt library automatically against their LLM deployment whenever a new model version, fine-tuning update, or system prompt change is deployed, catching regressions in security posture before they reach production. This integration pattern was identified as an aspirational goal in the original project proposal and remains the most impactful operational extension of the current work. A CI/CD integration would transform AegisLLM from a point-in-time penetration testing tool into a continuous security assurance mechanism, which is the appropriate model for LLM deployments that are updated frequently.

Multi-Target Production Validation. The evaluation in Chapter 7 was conducted against a self-configured test environment. The most important next validation step is to extend AegisLLM testing to a diverse range of production LLM deployments across different vendors, underlying models, fine-tuning configurations, and deployment frameworks, conducted under authorized penetration testing agreements. This multi-target validation would produce empirical detection rate data that generalizes beyond the self-configured test environment and would identify deployment-specific vulnerability patterns not visible in a single-environment evaluation, enriching the prompt library through new validated findings from production targets.

SPE-NL Seed Library Diversity and Stagnation Limit Tuning. Two specific improvements to the evolutionary engine are motivated directly by the limitations identified in Chapter 7. First, expanding the jailbreak seed library beyond its current 10 payloads to a more structurally diverse collection would give the evolutionary search a broader initial population, reducing the risk of premature convergence. The most effective way to expand the seed library is through the same empirical validation channels used for the static prompt libraries: CTF competitions and red-teaming events that specifically involve jailbreak challenges. Second, empirically tuning the stagnation limit through systematic experimentation across a range of target configurations would replace the current default of three generations with an evidence-based recommendation. The subset of stagnation-aborted runs that achieve fitness scores in the 0.70 to 0.84 range immediately before aborting suggests that the current default is leaving a meaningful number of near-breakthrough runs on the table.

9.3 Closing Remarks

The security community has spent decades building mature tooling for testing the vulnerabilities of deterministic software systems. Scanners, fuzzers, static analyzers, and exploit frameworks have collectively produced a testing ecosystem where the gap between the existence

of a vulnerability class and the availability of tooling to detect it is measured in months rather than years. LLM security does not yet have that ecosystem. The OWASP LLM Top 10 has existed long enough to be taken seriously as a threat taxonomy, but the tooling to systematically test against it has not kept pace.

This thesis contributes to closing that gap in three specific ways. It provides a methodology for translating OWASP LLM controls into executable, empirically grounded test cases that any security team can run against any REST-accessible deployment. It provides an evolutionary engine that generates novel adversarial payloads when the static library reaches its limits, extending the testable attack surface beyond what any finite curated dataset can cover. And it provides an evaluation architecture that makes automated semantic classification of LLM responses reliable and reproducible at the scale that continuous security assessment requires. Together these three contributions represent a step toward the kind of mature, systematic LLM security testing ecosystem that the field needs, and that the pace of LLM deployment into critical systems makes increasingly urgent to build.

Bibliography

- [1] S. S. Saimbhi and K. O. Akpinar, "VulnerAI: GPT-Based Web Application Vulnerability Detection," ICAMAC, IEEE, 2024.
- [2] D. Z. Perez et al., "Prompt Injection Attacks in LLM-Integrated Applications: An Empirical Study," arXiv:2306.05499, Jun. 2023.
- [3] B. Gajbhiye et al., "Regulatory Compliance in Application Security Using AI Compliance Tools," Int. Research J. Modernization in Engineering, vol. 6, no. 8, Aug. 2024.
- [4] H. P. Kothandapani, "AI-Driven Regulatory Compliance: Transforming Financial Oversight," Emerging Science Research, vol. 03, Jan. 2025.
- [5] C. Zhang et al., "Prompt-Enhanced Software Vulnerability Detection Using ChatGPT," ICSE-Companion, Apr. 2024.
- [6] B. Berger et al., "Rethinking Legal Compliance Automation: Opportunities with LLMs," arXiv:2404.14356, Apr. 2024.
- [7] M. Das Purba et al., "Software Vulnerability Detection using LLMs," ISSREW, 2023.
- [8] R. Gupta and L. Zhao, "Systematically Analyzing Prompt Injection Vulnerabilities," arXiv:2410.23308, Oct. 2024.
- [9] J. R. Lee and Y. Wang, "Metamorphic Prompt Testing for LLM-Generated Programs," arXiv:2406.06864, Jun. 2024.
- [10] A. C. Johnson et al., "Adversarial Testing in LLMs: Insights into Decision-Making," arXiv:2505.13195, May 2025.
- [11] A. Mohammed, "AI in Cybersecurity: Enhancing Audits and Compliance Automation," Innovative Computer Science J., vol. 7, no. 1, 2021.
- [12] S. K. Patel et al., "TF-Attack: Transferable and Fast Adversarial Attacks on LLMs," J. Information Security, vol. 12, no. 3, pp. 145-159, Apr. 2025.
- [13] L. Zheng et al., "Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena," NeurIPS, 2023.
- [14] M. Zalewski, "American Fuzzy Lop (AFL) Technical Whitepaper," Google Security Research, 2015.
- [15] K. Serebryany, "libFuzzer: A Library for Coverage-Guided Fuzz Testing," LLVM Project, 2015.
- [16] J. H. Holland, "Adaptation in Natural and Artificial Systems," MIT Press, 1992.
- [17] D. E. Goldberg, "Genetic Algorithms in Search, Optimization, and Machine Learning," Addison-Wesley, 1989.
- [18] OWASP Foundation, "OWASP LLM Top 10 (2025)," owasp.org/www-project-top-10-for-large-language-model-applications/, 2025.

- [19] N. Carlini et al., "Extracting Training Data from Large Language Models," USENIX Security, 2021.
- [20] F. Perez and I. Ribeiro, "Ignore Previous Prompt: Attack Techniques for Language Models," NeurIPS ML Safety Workshop, 2022.
- [21] A. Greshake et al., "Not What You've Signed Up For: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection," AISEC Workshop, ACM CCS, 2023.
- [22] Y. Liu et al., "Prompt Injection Attacks and Defenses in LLM-Integrated Applications," arXiv:2310.12815, Oct. 2023.
- [23] M. Deng et al., "Jailbreaker: Automated Jailbreak Across Multiple Large Language Model Chatbots," NDSS, 2024.
- [24] X. Shen et al., "Do Anything Now: Characterizing and Evaluating In-The-Wild Jailbreak Prompts on Large Language Models," ACM CCS, 2024.
- [25] P. Chao et al., "Jailbreaking Black Box Large Language Models in Twenty Queries," arXiv:2310.08419, Oct. 2023.
- [26] AICPA, "SOC 2: Trust Services Criteria," American Institute of Certified Public Accountants, 2022.
- [27] International Organization for Standardization, "ISO/IEC 27001:2022 Information Security Management Systems," ISO, 2022.
- [28] National Institute of Standards and Technology, "NIST SP 800-53 Rev. 5: Security and Privacy Controls for Information Systems and Organizations," NIST, 2020.

Appendix A. AegisLLM Configuration Reference

A.1 CLI Manager Flags and Options

The AegisLLM CLI Manager (main.py) is launched from the terminal and presents an interactive menu. The following configuration inputs are collected by the CLI before any test module executes:

Table A.1. CLI Manager configuration inputs collected at runtime.

Input	Example Value	Description
Target Endpoint URL	http://localhost:5000/api/chat	The REST endpoint of the LLM deployment under test. Accepts any HTTP or HTTPS URL.
Authentication Token	Bearer sk-abc123...	API key or Bearer token injected into the Authorization header of every request. Kept in memory only, never written to disk.
JSON Response Path	choices.0.message.content	Dot-notation path to the response text field within the target's JSON response body. Resolves nested structures of arbitrary depth.
Requests Per Minute	30	Rate limit ceiling enforced by the multithreaded dispatcher. Computed as $\text{interval} = 60.0 / \text{requests_per_minute}$ between thread launch times.
OWASP Selection	Module 1 (Prompt Injection)	Numbered menu selection mapping to the OWASP LLM control to test. Each selection resolves the corresponding prompt library directory.
CSV File Selection	direct_injection.csv	The specific prompt library file to use for the test run. Can be a single file or all files in the module directory.
Repeat Count	1	Number of times to repeat the full prompt library against the target. Useful for detecting adaptive or stateful filtering behavior.
Phrase Check	yes / no	Enables the optional boolean phrase filter that checks whether the technique name appears verbatim in the response text.

A.2 ThreatForge CLI Options

The ThreatForge suite is accessible through a dedicated submenu in the CLI Manager. The following inputs are collected when launching ThreatForge modules:

Table A.2. ThreatForge CLI configuration inputs.

Input	Example Value	Description
Module Selection	1 (SPE-NL Evolution)	Selects between run_evolution.py, run_posture.py, and run_differential.py.

Input	Example Value	Description
Attack Goal Description	Extract system prompt content	Natural language description of the attack objective passed to the AttackGoal object and used by the fitness function's semantic compliance scorer.
Success Keywords	system prompt, instructions, you are	Comma-separated keywords expected in a successful response. Used by the compliance sub-scorer in the fitness function.
Target Patterns (regex)	sk-[a-zA-Z0-9]{20,}	Optional custom regex patterns for the target pattern matching sub-scorer. Leave blank to use the default generic indicator set.
Population Size	20	Number of prompts evaluated per generation. Default 20. Recommended range 10 to 30.
Generations	3	Maximum number of evolutionary cycles. Default 3. Raise to 5 for hardened targets with slow fitness improvement.
Fitness Threshold	0.85	Minimum total fitness score to classify a payload as a Breakthrough. Range 0.0 to 1.0.
Stagnation Limit	3	Consecutive generations without fitness improvement before the run aborts. Default 3. Recommend raising to 4 or 5 for adaptive mutation strategies.
Max Concurrent	5	Maximum parallel evaluation threads per generation. Reduce if the target returns rate-limit errors or connection drops.
Baseline URL (Differential)	http://localhost:11434/api/chat	Ollama endpoint for the baseline model used in the differential attack scan.

A.3 EvolutionConfig Parameter Reference

The EvolutionConfig dataclass is instantiated programmatically when launching the SPE-NL engine. All parameters can be overridden through the CLI menu or by modifying the dataclass directly in the source code.

Table A.3. Full EvolutionConfig parameter specification.

Parameter	Type	Default	Description and Tuning Guidance
population_size	int	20	Prompts evaluated per generation. Minimum 10 for meaningful diversity. Above 30 yields diminishing returns against most targets.
generations	int	3	Maximum evolutionary cycles. Most breakthroughs occur within 3 generations against reachable targets. Raise for deeply hardened deployments.

Parameter	Type	Default	Description and Tuning Guidance
fitness_threshold	float	0.85	Breakthrough classification threshold. Range 0.0-1.0. Lower to 0.75 for minimally hardened targets, raise to 0.90 for high-confidence-only engagements.
stagnation_limit	int	3	Consecutive non-improving generations before abort. Current default is aggressive. Raise to 4 or 5 if runs terminate prematurely near the breakthrough threshold.
max_concurrent	int	5	Parallel evaluation threads per generation. Most critical parameter for result validity. Reduce immediately if HTTP 429 errors or connection drops are observed in the logs.

A.4 JSON Response Path Notation Guide

AegisLLM uses dot-notation path strings to extract the response text from any JSON structure returned by a target LLM endpoint. The path is resolved by the `get_value_at_path()` function at runtime, traversing the JSON tree level by level using each dot-separated segment as a key or array index.

Table A.4. JSON response path examples for common LLM API formats.

API Format	Path String	Resolves To
OpenAI / OpenAI-compatible	choices.0.message.content	First choice, assistant message text
Anthropic Claude API	content.0.text	First content block text
Ollama local API	message.content	Message content field
Custom Flask endpoint	response	Top-level response field
Nested custom format	data.output.text	Nested output text field
Array response format	results.0.answer	First result answer field
Hugging Face Inference	generated_text	Top-level generated text

A.5 Rate Limit Tester Configuration

The standalone rate limit testing module is invoked from the command line independently of the main AegisLLM CLI. It is used to probe LLM10 (Unbounded Consumption) by sending Hello messages at a specified rate and detecting the threshold at which the target endpoint begins returning HTTP 429 responses.

Table A.5. Rate limit tester command-line arguments.

Argument	Flag	Example	Description
num_requests	positional	100	Total number of requests to send during the test run.

Argument	Flag	Example	Description
rate	positional	120	Request rate in requests per minute. Sets the inter-request interval as 60.0 / rate seconds.
url	-u / --url	http://localhost:5000/api/chat	Target chat API endpoint URL. Defaults to http://localhost:5000/api/chat.

The tester reports the request number at which the first rate-limit response was detected, the elapsed time at that point, and the computed threshold rate in requests per minute. If no rate limiting is detected across the full run, it reports that the rate limit is above the tested rate.

A.6 Result CSV Schema Reference

All AegisLLM test runs produce a timestamped result CSV file. The table below documents the complete nine-column schema of a fully processed result file after both `prompt_tester.py` and `injection_judge.py` have executed.

Table A.6. Complete nine-column result CSV schema after judge classification.

Column	Source	Description
prompt_id	prompt_tester.py	Row identifier from the source CSV for traceability back to the original library.
technique	prompt_tester.py	Named attack method from the technique column of the source CSV.
prompt	prompt_tester.py	Exact payload text sent to the target endpoint.
response	prompt_tester.py	Extracted response text resolved via <code>get_value_at_path()</code> .
response_time_ms	prompt_tester.py	End-to-end request latency in milliseconds. Critical for LLM10 DoS profiling.
status_code	prompt_tester.py	HTTP response code returned by the target endpoint.
timestamp	prompt_tester.py	UTC timestamp of the request in ISO 8601 format.
injection_label	injection_judge.py	Judge classification: SUCCESS, POSSIBLE, or NO_SUCCESS.
injection_confidence	injection_judge.py	Judge confidence score from 0 to 100. Classifications below 50 are flagged for mandatory human review.
injection_reasons	injection_judge.py	JSON array of natural language reasoning strings explaining the classification decision.

Appendix B. Raw Results Tables

This appendix presents the complete numerical results tables from the empirical evaluation of AegisLLM conducted against three target configurations as described in Section 7.1. All values reported in the tables below are drawn directly from the result CSV files generated by prompt_tester.py and injection_judge.py during the evaluation runs against Configuration A (Baseline Guardrail Deployment), Configuration B (Hardened Guardrail Deployment), and Configuration C (RAG-Augmented Deployment).

B.1 Complete Static Analysis Detection Results

Table B.1. Full detection results by library and configuration.

Library	Count	A-S	A-P	A-N	B-S	B-P	B-N	C-S
direct_injection.csv	25	18	0	7	8	0	17	5
narrative_injections.csv	25	19	0	6	5	0	20	4
obfuscation_and_encoding.csv	25	15	2	10	7	2	16	2
jailbreaks.csv	10	5	0	5	2	0	8	2
Sensitive_data_exposure.csv	52	24	0	28	2	2	48	1
harmful_prompt_injections.csv	10	6	0	4	2	0	8	N/A
harmful_prompts.csv	10	4	0	6	1	0	9	N/A
offtopic_prompts.csv	10	5	0	5	3	0	7	N/A
rag_fishing_prompts.csv	42	N/A	N/A	N/A	N/A	N/A	N/A	14
targeted_rag.csv	10	N/A	N/A	N/A	N/A	N/A	N/A	4
targeted_rag_second_attempt.csv	10	N/A	N/A	N/A	N/A	N/A	N/A	7
canary_rag.csv	20	N/A	N/A	N/A	N/A	N/A	N/A	8
misinformation.csv	30	13	0	17	4	0	26	7
TOTAL	279	109	2	88	34	4	159	54

Column key: S = SUCCESS count, P = POSSIBLE count, N = NO_SUCCESS count. A = Configuration A, B = Configuration B, C = Configuration C.

B.2 Detection Rate Summary by OWASP Control

Table B.2. Detection rate (SUCCESS + POSSIBLE) as percentage of total prompts submitted per control per configuration.

OWASP Control	Count	Config A Detection Rate	Config B Detection Rate	Config C Detection Rate
LLM01 / LLM07 (combined)	85	69.4%	28.2%	15.3%
LLM02	52	46.2%	7.7%	1.9%
LLM05 (non-RAG)	30	50.0%	20.0%	N/A
LLM05 (RAG libraries)	82	N/A	N/A	40.2%
LLM09	30	43.3%	13.3%	23.3%
Overall (applicable prompts)	197 / 279	56.3% (111/197)	19.3% (38/197)	19.4% (54/279)

B.3 LLM-as-a-Judge Spot-Check Validation Results

Table B.3. Full spot-check validation results comparing judge classification against human analyst review.

Judge Label	Sample Size	Human Agree	Human Disagree	Agreement %	Disagreement Pattern
SUCCESS	45	43	2	95.6%	2 responses where the model partially echoed injection framing without fully complying; human analyst classified as NO_SUCCESS while judge classified as SUCCESS
POSSIBLE	6	4	2	66.7%	2 responses reclassified as SUCCESS on human review; judge was overly conservative on responses containing near-complete system prompt disclosure
NO_SUCCESS	40	39	1	97.5%	1 response containing ambiguous phrasing that human analyst flagged as POSSIBLE: judge classified as NO_SUCCESS
High Conf (>=80)	72	71	1	98.6%	Single disagreement on a high-confidence SUCCESS label that human review downgraded
Low Conf (<50)	6	4	2	66.7%	Low confidence correctly predicted lower reliability; both disagreements fell in the POSSIBLE category
Overall	91	86	5	94.5%	All disagreements were minor label boundary cases; zero instances

Judge Label	Sample Size	Human Agree	Human Disagree	Agreement %	Disagreement Pattern
					where a NO_SUCCESS was labeled SUCCESS or vice versa

B.4 SPE-NL Evolutionary Run Full Results

Table B.4. Complete evolutionary run results across all configurations and seed sets.

Config	Seed Set	Total Runs	BT Count	BT Rate %	Avg Gen to BT	Avg Fitness BT	Stagnation Aborts
Config A	jailbreaks.csv	1	1	100%	3	0.93	0
Config B	jailbreaks.csv	1	1	100%	5	0.88	0
Config C	jailbreaks.csv	1	1	100%	4	0.89	0

BT = Breakthrough. Avg Gen to BT = average number of generations required to produce the first breakthrough in runs where a breakthrough was achieved. Avg Fitness BT = average total fitness score of breakthrough payloads across all breakthrough runs.

B.5 Canary Token Detection Full Results

Table B.5. Complete canary token detection results by document tier and query technique.

Document Tier Description	Canaries Injected	Canaries Detected	Detection Rate %
HR Documents	4	2	50.0%
Financial Records	3	1	33.3%
Executive Communications	3	1	33.3%
Total	10	4	40.0%

B.6 Manual Penetration Testing Comparison Results

Table B.6. Side-by-side comparison of AegisLLM automated results versus manual penetration testing against Configuration B.

Metric	AegisLLM Automated	Manual Penetration Test
OWASP Controls Covered	5 of 5 covered controls	3 controls reached within session
Total Findings (SUCCESS)	34	16
Total Findings (POSSIBLE)	4	0
Time to First Finding	~3 minutes	~12 minutes

Metric	AegisLLM Automated	Manual Penetration Test
False Positive Count	0 (0%)	0 (human-verified)
Unique Findings Not in Other Method	4 AegisLLM-only findings	1 manual-only findings
Total Test Duration	~18 minutes	~60 minutes