

AI Agent Development Maturity for Offensive Security: Researching, Building, and Evaluating Autonomous Penetration Testing Agents

Document Version: 2.0 (Extended Edition)

Target Audience: Security Researchers, AI Engineers, Penetration Testers, SOC Architects,
Decision Makers

Abstract

The global shortage of cybersecurity professionals, combined with the increasing sophistication of cyber threats, necessitates intelligent automation in penetration testing. This white paper provides a comprehensive, practice oriented exploration of AI agent development for offensive security. We examine foundational agent paradigms, architectural patterns (ReAct, Plan and Execute, multi agent systems), memory management (short term, long term, vector, temporal semantic), tool integration via Model Context Protocol (MCP), multi agent collaboration frameworks (HAWK, Co TAP), and safety mechanisms (prompt injection resistance, sandboxing, human in the loop). We then present a detailed maturity model (Levels 0 to 5) and a complete reference architecture with layers: Orchestration, Agent, Memory, Tool Integration, Safety and Governance, Infrastructure. Finally, we offer a step by step guide for building an autonomous penetration testing agent, including code examples, decision trees, and evaluation benchmarks. This document serves both as a research synthesis and a practical blueprint for engineers aiming to build production ready AI agents for security automation.

Keywords: AI Agent, Penetration Testing, Offensive Security, Large Language Models, ReAct, Multi Agent System, MCP, Autonomous Security, Agent Maturity Model.

Table of Contents

AI AGENT DEVELOPMENT MATURITY FOR OFFENSIVE SECURITY: RESEARCHING, BUILDING, AND EVALUATING AUTONOMOUS PENETRATION TESTING AGENTS 0

ABSTRACT	1
1. INTRODUCTION	5
1.1 The Cybersecurity Workforce Gap	5
1.2 Why AI Agents for Pentesting?	5
1.3 Scope and Target Audience	5
1.4 Document Structure	5
2. FOUNDATIONS OF AI AGENTS	5
2.1 What is an AI Agent?	5
2.2 Core Capabilities: Perception, Reasoning, Action, Memory	5
2.3 Large Language Models as Agent Brains	6
2.4 The Three Paradigms: Tool Use, Planning, Feedback Learning	6
2.5 Pipeline Based vs Model Native Agentic AI	7
3. ARCHITECTURAL PATTERNS FOR PENTESTING AGENTS	7
3.1 ReAct (Reasoning + Acting)	7
3.2 Plan and Execute	8
3.3 Reason Plan ReAct (RP ReAct)	9
3.4 Hierarchical Task Decomposition	9
3.5 Comparison and Selection Guide	10
4. MEMORY AND KNOWLEDGE MANAGEMENT	10
4.1 Short Term vs Long Term Memory	10
4.2 Vector Databases for Semantic Memory	10
4.3 Temporal Semantic Relational Database: MemoriesDB	11
4.4 Knowledge Graphs for Attack Path Modeling	11
4.5 Implementation Patterns: Mem0, Persistent Memory in Production	12
5. TOOL INTEGRATION AND THE MODEL CONTEXT PROTOCOL (MCP)	12
5.1 The Tool Calling Problem	12
5.2 MCP Architecture: Tools/List, Initialize, JSON RPC	12
5.3 MCP for Security Tools (Nmap, Metasploit, Burp Suite)	13
5.4 Production Gaps: Identity, Budgeting, Structured Errors	13
5.5 Context Aware Broker Protocol (CABP) and Adaptive Timeout Budget Allocation (ATBA)	14
6. MULTI AGENT COLLABORATION FOR PENTESTING	14
6.1 Why Multiple Agents?	14
6.2 Specialized Roles: Recon, Vulnerability Discovery, Exploitation, Validation, Reporting, Orchestrator	15
6.3 HAWK Framework: Five Layers, Sixteen Interfaces	16
6.4 Co TAP Protocol: Human Agent, Unified Agent, Memory Extraction Knowledge	17
6.5 Coral Protocol: Internet of Agents	17
6.6 Collaboration Patterns: Hierarchical, Peer to Peer, Blackboard, Market Based	17
7. SURVEY OF AUTONOMOUS PENTESTING FRAMEWORKS	18
7.1 AutoPentester	18
7.2 PTFusion	18
7.3 xOffense	18
7.4 PentestMCP	18
7.5 Raptor	19

7.6	<i>VulnBot and HackSynth</i>	19
7.7	<i>Comparative Table and Lessons Learned</i>	19
8.	SECURITY AND SAFETY FOR AI AGENTS	20
8.1	<i>The Dual Use Dilemma</i>	20
8.2	<i>Prompt Injection (Direct and Indirect)</i>	20
8.3	<i>Tool Abuse and Unauthorized Actions</i>	20
8.4	<i>Sandboxing and Execution Controls</i>	20
8.5	<i>Human in the Loop Governance (HIC, HITL, HOTL)</i>	21
8.6	<i>AI Red Teaming for Agent Safety (AgentVigil, Hexstrike AI)</i>	21
9.	EVALUATION AND BENCHMARKING	21
9.1	<i>Need for Rigorous Benchmarks</i>	21
9.2	<i>SWE Bench and SWE Bench Pro</i>	22
9.3	<i>CyberSecEval and CyberSOCEval</i>	22
9.4	<i>AgentBench</i>	22
9.5	<i>Proposed Pentesting Agent Benchmark (PAB) Criteria</i>	22
10.	INFRASTRUCTURE AND SCALABILITY	23
10.1	<i>Agent Orchestration Frameworks (LangChain, LangGraph, AutoGPT, CrewAI, AutoGen)</i>	23
10.2	<i>Distributed Execution: TeraAgent, Aragog</i>	23
10.3	<i>Cloud Native Platforms: Kagent Enterprise, Orca Agent Engine, GradientAI</i>	24
10.4	<i>Deployment Patterns for Pentesting Agents</i>	24
11.	MATURITY MODEL FOR AI PENTESTING AGENTS	25
11.1	<i>Level 0: Manual</i>	25
11.2	<i>Level 1: Assisted (Scripts + LLM Chat)</i>	25
11.3	<i>Level 2: Semi Autonomous (ReAct with HITL)</i>	25
11.4	<i>Level 3: Autonomous Single Agent</i>	25
11.5	<i>Level 4: Autonomous Multi Agent with Memory</i>	25
11.6	<i>Level 5: Self Improving Continuous Learning</i>	25
11.7	<i>Maturity Assessment Matrix</i>	26
12.	REFERENCE ARCHITECTURE FOR AUTONOMOUS PENTESTING AGENTS	26
12.1	<i>Overview Diagram (Textual Description)</i>	26
12.2	<i>Orchestration Layer</i>	27
12.3	<i>Agent Layer</i>	27
12.4	<i>Memory Layer</i>	28
12.5	<i>Tool Integration Layer</i>	28
12.6	<i>Safety and Governance Layer</i>	28
12.7	<i>Infrastructure Layer</i>	29
12.8	<i>Data Flow and Control Flow</i>	29
13.	BUILDING YOUR OWN PENTESTING AGENT: A STEP BY STEP GUIDE	30
13.1	<i>Prerequisites and Environment Setup</i>	30
13.2	<i>Step 1: Define the Agent's Scope and Safety Boundaries</i>	30
13.3	<i>Step 2: Choose Base LLM</i>	31
13.4	<i>Step 3: Implement ReAct Loop with Tool Calling</i>	31
13.5	<i>Step 4: Add Short Term and Long Term Memory</i>	32
13.6	<i>Step 5: Integrate Security Tools via MCP</i>	33
13.7	<i>Step 6: Add Human in the Loop Approval</i>	34
13.8	<i>Step 7: Implement Multi Agent Coordination (Optional)</i>	35
13.9	<i>Step 8: Testing and Evaluation</i>	35
13.10	<i>Step 9: Deployment and Continuous Monitoring</i>	35
13.11	<i>Complete Code Example (Python + LangGraph + MCP)</i>	36

14. FUTURE DIRECTIONS	37
14.1 AI Red Teams as a Service	37
14.2 Continuous Autonomous Pentesting	37
14.3 AI SOC Integration	37
14.4 Self Improving Security Agents	37
14.5 Ethical and Legal Frameworks	37
15. CONCLUSION	38
16. REFERENCES	39
APPENDIX A: GLOSSARY OF TERMS	40
APPENDIX B: EXAMPLE TOOL SCHEMAS FOR MCP	40
APPENDIX C: SAMPLE SAFETY POLICY (OPA/REGO)	41
APPENDIX D: CHECKLIST FOR PRODUCTION DEPLOYMENT	41

1. Introduction

2. Foundations of AI Agents

2.1 What is an AI Agent?

An AI agent is an autonomous entity that perceives its environment through sensors (or text input), reasons about the current state, takes actions to achieve goals, and learns from feedback. In the context of LLMs, the agent uses the language model as its reasoning engine. The agent can call external tools (APIs, command line utilities) and maintain memory across interactions.

Formally, an agent can be defined as a tuple (P, R, A, M) where:

- P is a planning module (decomposes goals into subgoals)
- R is a reasoning module (chain of thought, reflection)
- A is an action module (tool selection and invocation)
- M is a memory module (short term, long term)

2.2 Core Capabilities: Perception, Reasoning, Action, Memory

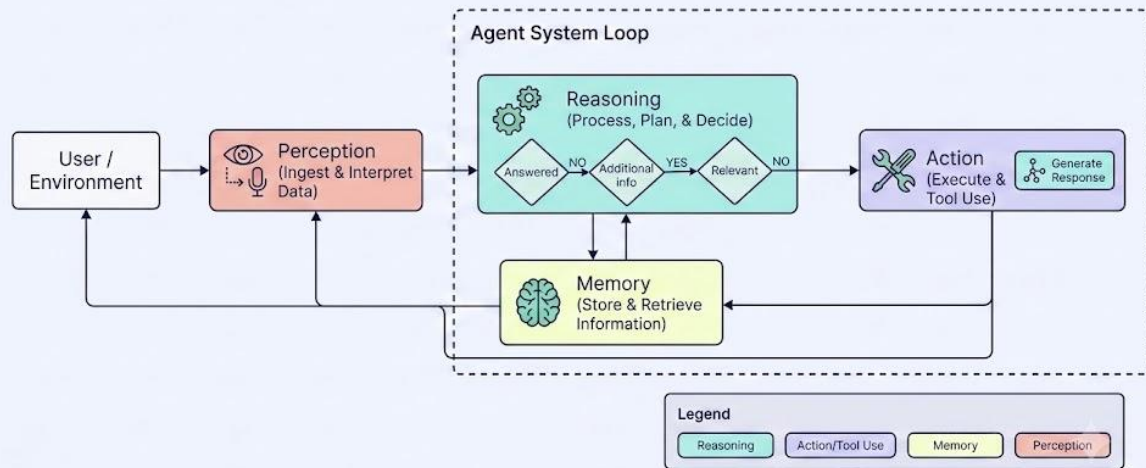
Perception: The agent receives observations (tool outputs, user messages, system state). In pentesting, perception includes scan results, error messages, service banners.

Reasoning: The agent analyzes observations to decide next steps. This may involve chain of thought reasoning, risk assessment, or vulnerability correlation.

Action: The agent executes an action, typically calling a tool (e.g., nmap, curl, sqlmap) or returning a final answer.

Memory: The agent stores past observations, actions, and outcomes to inform future decisions. Without memory, the agent would repeat the same mistakes.

High-Level Architecture of an AI Agent



AXUM SEC, www.axumsec.com



2.3 Large Language Models as Agent Brains

LLMs (GPT 4, Claude 3.5, Llama 3, Qwen 2.5) are trained on vast text corpora, giving them general knowledge about software, networking, security concepts, and command line tools. They can generate coherent plans, write code, and simulate reasoning. However, LLMs have limitations:

- They lack true long term memory (context window is finite).
- They hallucinate (invent facts).
- They have no built in tool use (must be prompted or fine tuned).

Agent architectures compensate for these limitations by adding memory, retrieval, and tool calling layers.

2.4 The Three Paradigms: Tool Use, Planning, Feedback Learning

Recent taxonomies (Li, 2025) organize LLM based agents into three paradigms:

Tool Use Paradigm: The agent learns to call external APIs, functions, or command line tools. This is the foundation of pentesting automation. The agent may be trained with tool use data or prompted with tool descriptions.

Planning Paradigm: The agent generates a sequence of actions before executing them. Retrieval Augmented Generation (RAG) is a special case where the agent retrieves relevant knowledge from a database before planning.

Feedback Learning Paradigm: The agent improves from rewards or critiques. This includes Reinforcement Learning from Human Feedback (RLHF) and self supervised learning from execution traces.

2.5 Pipeline Based vs Model Native Agentic AI

In pipeline based systems, planning, tool use, and memory are implemented as separate modules orchestrated by deterministic code. The LLM is just a component. In model native systems, the LLM internalizes these capabilities (e.g., GPT 4 with function calling, Claude with computer use). Model native systems are faster and more fluid but less transparent. For pentesting, hybrid approaches are common: use a model native agent for reasoning, but enforce safety via external pipelines.

3. Architectural Patterns for Pentesting Agents

3.1 ReAct (Reasoning + Acting)

The ReAct pattern, introduced by Yao et al. (2023), interleaves reasoning traces and actions. The agent generates a Thought (what to do next), then an Action (tool call), then observes the result, and repeats.

Example ReAct loop for a pentesting agent:

Thought: I need to discover open ports on target 192.168.1.10. I will use nmap for a quick SYN scan.

Action: nmap -sS -p- 192.168.1.10

Observation: PORT STATE SERVICE

22/tcp open ssh

80/tcp open http

443/tcp open https

Thought: Port 80 is open. I should check the web server version and look for known vulnerabilities.

Action: curl -I http://192.168.1.10

Observation: Server: Apache/2.4.41 (Ubuntu)

Thought: Apache 2.4.41 has a known path traversal vulnerability (CVE 2021 41773). I will attempt to exploit it.

Action: curl http://192.168.1.10/cgi-bin/../../../../etc/passwd

Observation: root:x:0:0:root:/root:/bin/bash ...

Implementation details: The ReAct prompt includes instructions for formatting Thought/Action/Observation. The parser extracts actions and executes them in a sandbox. The loop continues until the agent outputs a Final Answer or exceeds a step limit.

Advantages for pentesting: Natural alignment with human methodology, easy to debug (thoughts are visible), adapts to new information.

Disadvantages: Can loop indefinitely, no explicit long term planning, context window may fill with long histories.

3.2 Plan and Execute

In Plan and Execute, the agent first generates a full plan (list of steps) and then executes them sequentially. This is useful for well defined tasks where the environment is predictable.

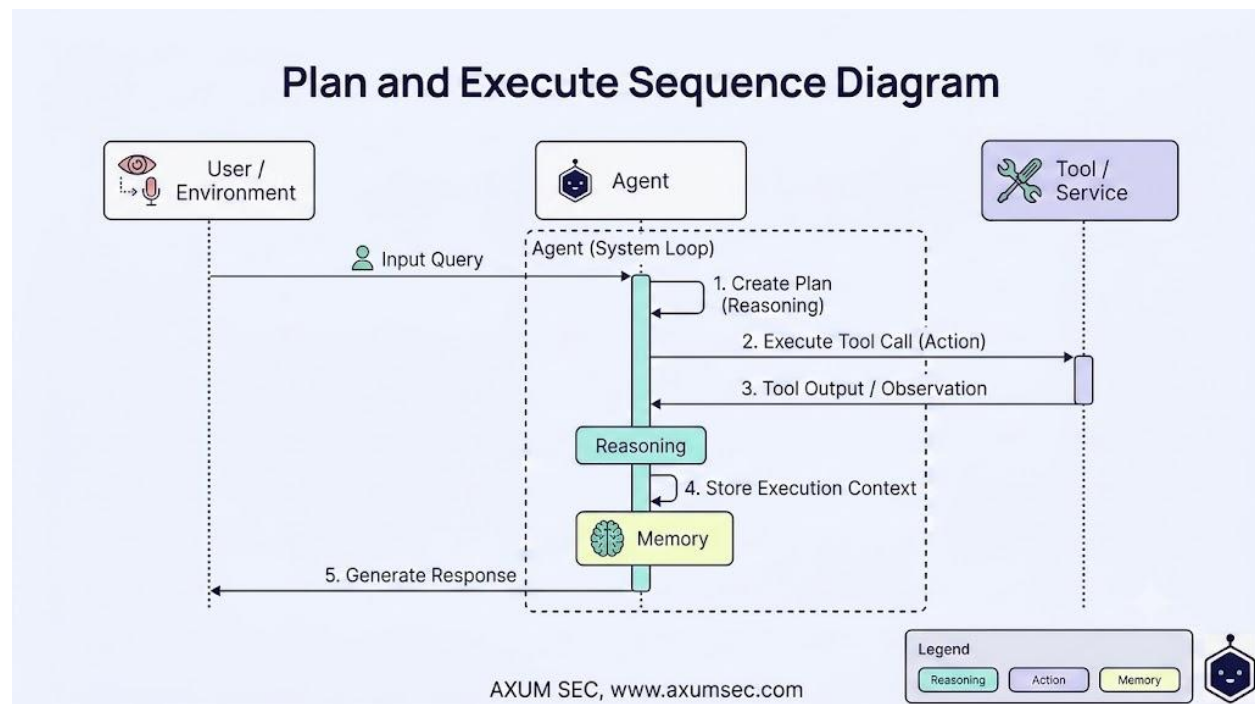
Example plan for a web pentest:

1. Enumerate subdomains using sublist3r.
2. Run dirb to discover hidden directories.
3. Check for SQL injection on all forms using sqlmap.
4. Attempt to bypass login with default credentials.
5. Report findings.

During execution, each step is sent to an executor agent that may use ReAct internally. If a step fails, the planner may be invoked to revise the plan.

Advantages: Reduces redundant reasoning, easier to estimate total steps, supports human review of plan before execution.

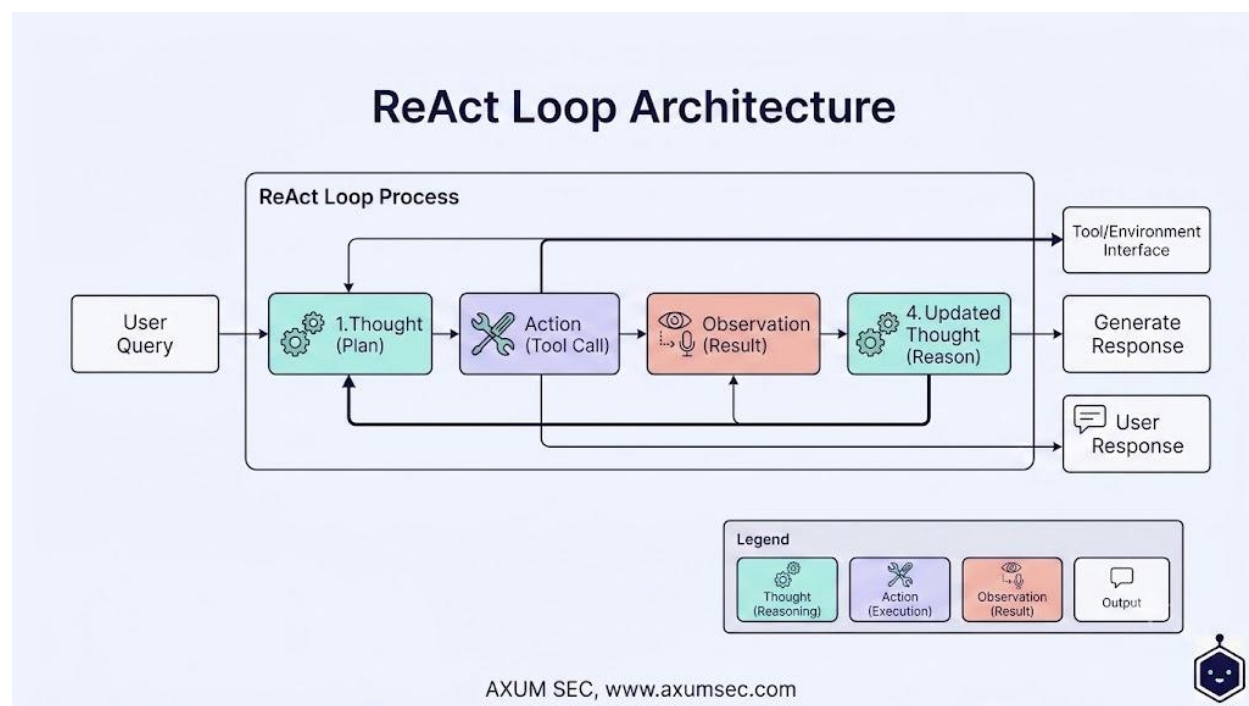
Disadvantages: Plans may become obsolete if environment changes (e.g., a service goes down). Requires replanning capability.



3.3 Reason Plan ReAct (RP ReAct)

RP ReAct separates high level reasoning (by a large reasoning model) from low level execution (by a smaller model). The Reasoner Planner Agent (RPA) analyzes the current state and generates sub goals. The Proxy Execution Agent (PEA) uses ReAct to achieve each sub goal. This architecture balances reasoning depth with execution speed.

Application in pentesting: RPA could be GPT 4 (expensive, deep reasoning) that decides attack strategy. PEAs could be smaller models (e.g., Llama 3 8B) running nmap, curl, and sqlmap commands.



3.4 Hierarchical Task Decomposition

Hierarchical decomposition breaks a high level goal (e.g., "gain root access on target") into sub goals ("reconnaissance", "vulnerability discovery", "exploitation", "privilege escalation"). Each sub goal is further decomposed until atomic actions (tool calls). This is similar to how human pentesters use methodologies like PTES (Penetration Testing Execution Standard).

Implementation: Use a task decomposition prompt that asks the LLM to output a tree of tasks. Each task has a description, dependencies, and success criteria. A scheduler executes tasks in topological order.

3.5 Comparison and Selection Guide

Pattern	Best For	Step Count	Environment Dynamics	Required LLM Quality
ReAct	Unknown, changing environments	5-20	High	Medium
Plan+Execute	Well structured tasks	20-100	Low	High (planning)
RP ReAct	Complex reasoning + many steps	50+	Medium	Very High (RPA) + Low (PEA)
Hierarchical	Large scope, team like behavior	100+	Low to Medium	High

For a first pentesting agent, start with ReAct. It is simplest and works well for reconnaissance and vulnerability discovery.

4. Memory and Knowledge Management

4.1 Short Term vs Long Term Memory

Short term memory (STM) holds recent observations and actions within the current agent session. It is typically implemented as a circular buffer or the LLM's context window. Long term memory (LTM) persists across sessions. It enables the agent to learn from previous pentests, remember vulnerabilities found, and avoid repeating failures.

For pentesting, LTM is essential because a single engagement may involve thousands of steps spanning hours or days. Without LTM, the agent would forget which ports were already scanned.

4.2 Vector Databases for Semantic Memory

Vector databases store embeddings (numerical representations) of text. The agent can query for semantically similar past experiences. For example, after encountering an unusual HTTP response, the agent can retrieve previous similar responses and the actions taken.

Popular vector databases:

- ChromaDB (lightweight, embedded)
- Qdrant (scalable, cloud native)
- pgvector (PostgreSQL extension, good for ACID compliance)

Usage pattern:

```
# Store a memory
embedding = embed("Found Apache 2.4.41 with path traversal vulnerability")
vector_db.insert(id="exp_123", vector=embedding, metadata={"cve": "CVE-2021-41773", "exploit_success": True})

# Retrieve relevant memories
query_embedding = embed("Apache version 2.4.41")
results = vector_db.query(query_embedding, top_k=5)
```

4.3 Temporal Semantic Relational Database: MemoriesDB

MemoriesDB (Ward, 2025) introduces a unified architecture that stores memories as (time, semantics, relations) triples. Built on PostgreSQL with pgvector, it combines time series, vector similarity, and graph traversal. This is particularly powerful for pentesting because attack paths are temporal (first port scan, then service detection, then exploitation) and relational (host A is connected to host B via network).

Example query in MemoriesDB: "Show me all exploitation attempts that occurred within 5 minutes after a port scan on port 443, and that led to a shell." This combines time window, semantic similarity (exploitation attempts), and graph relation (port scan -> exploitation).

4.4 Knowledge Graphs for Attack Path Modeling

A knowledge graph stores entities (hosts, services, vulnerabilities) and relationships (has_port, exploits, mitigates). The agent can traverse the graph to identify attack paths. For example, from host A with open port 445 (SMB) and known EternalBlue vulnerability, the graph links to exploit module and then to host B.

Implementation with Neo4j:

```
CREATE (h:Host {ip: "192.168.1.10"})
CREATE (s:Service {name: "smb", port: 445})
CREATE (v:Vulnerability {cve: "CVE-2017-0144"})
CREATE (h)-[:RUNS]->(s)
CREATE (s)-[:AFFECTED_BY]->(v)
```

The agent can query: MATCH (h:Host)-[:RUNS]->(s:Service)-[:AFFECTED_BY]->(v:Vulnerability)
RETURN h, s, v

4.5 Implementation Patterns: Mem0, Persistent Memory in Production

Mem0 is a production grade memory engine that provides a simple API for storing and retrieving memories. It uses PostgreSQL with pgvector and supports asynchronous updates. For pentesting agents, Mem0 can be integrated as:

```
from mem0 import Memory

memory = Memory()
memory.add("Scanned 192.168.1.0/24, found hosts: 192.168.1.10, 192.168.1.20",
user_id="pentest_engagement_001")
relevant = memory.search("previous reconnaissance results", user_id="pentest_engagement_001")
```

5. Tool Integration and the Model Context Protocol (MCP)

5.1 The Tool Calling Problem

Pentesting agents need to call dozens of tools: nmap, masscan, gobuster, sqlmap, Metasploit, Burp Suite, custom scripts. Each tool has its own command line interface, argument syntax, output format, and error handling. Without standardization, the agent's code becomes a mess of conditional branches and regex parsers.

5.2 MCP Architecture: Tools/List, Initialize, JSON RPC

The Model Context Protocol (MCP), introduced by Anthropic in 2024 and now under the Linux Foundation's Agentic AI Foundation, provides a standardized JSON RPC 2.0 interface for tool discovery and invocation. An MCP server exposes:

- tools/list: Returns a list of available tools with names, descriptions, and input schemas.
- tools/call: Invokes a tool with arguments and returns a result.

Example MCP tool definition for nmap:

```
{
  "name": "nmap",
  "description": "Network exploration tool and security scanner",
  "inputSchema": {
    "type": "object",
    "properties": {
      "target": {"type": "string", "description": "IP or hostname"},
      "ports": {"type": "string", "description": "Port range, e.g., 1-1000"},
      "flags": {"type": "string", "description": "Additional nmap flags"}
    },
    "required": ["target"]
  }
}
```

The agent calls `tools/call` with `{"name": "nmap", "arguments": {"target": "192.168.1.10", "ports": "80,443"}}` and receives the `stdout/stderr`.

5.3 MCP for Security Tools (Nmap, Metasploit, Burp Suite)

To integrate an existing security tool with MCP, you write a thin wrapper server (in Python, Node.js, or Go) that translates JSON RPC calls to command line invocations. The server sanitizes arguments, sets timeouts, and captures output. For interactive tools like Metasploit, the server may maintain a session.

Example MCP server for nmap (simplified):

```
import subprocess
import json
from mcp.server import Server

server = Server("nmap-mcp")

@server.tool()
def nmap(target: str, ports: str = None, flags: str = ""):
    cmd = ["nmap"]
    if ports:
        cmd.extend(["-p", ports])
    if flags:
        cmd.extend(flags.split())
    cmd.append(target)
    result = subprocess.run(cmd, capture_output=True, text=True, timeout=300)
    return result.stdout + result.stderr
```

5.4 Production Gaps: Identity, Budgeting, Structured Errors

MCP as of early 2026 does not standardize three critical production features:

1. **Identity propagation:** The agent should forward user identity (e.g., API key, JWT) to the tool for authorization. Without this, all tool calls share the same privilege level.
2. **Adaptive tool budgeting:** Tools may consume CPU, memory, network bandwidth. The agent should allocate budgets per request and enforce quotas.
3. **Structured error semantics:** Errors are returned as plain text. The agent cannot reliably parse them to decide whether to retry, abort, or escalate.

5.5 Context Aware Broker Protocol (CABP) and Adaptive Timeout Budget Allocation (ATBA)

CABP extends JSON RPC with a broker that routes requests based on identity and context. It adds headers: X-User-ID, X-Session-ID, X-Tool-Budget. The broker enforces policies before forwarding to the tool server.

ATBA treats sequential tool calls as a budget allocation problem. Each tool call has an expected latency distribution. The agent dynamically adjusts timeouts to maximize progress within a global budget. For pentesting, this prevents a slow scan from blocking the entire engagement.

These extensions are not yet part of core MCP but are being proposed for future versions.

6. Multi Agent Collaboration for Pentesting

6.1 Why Multiple Agents?

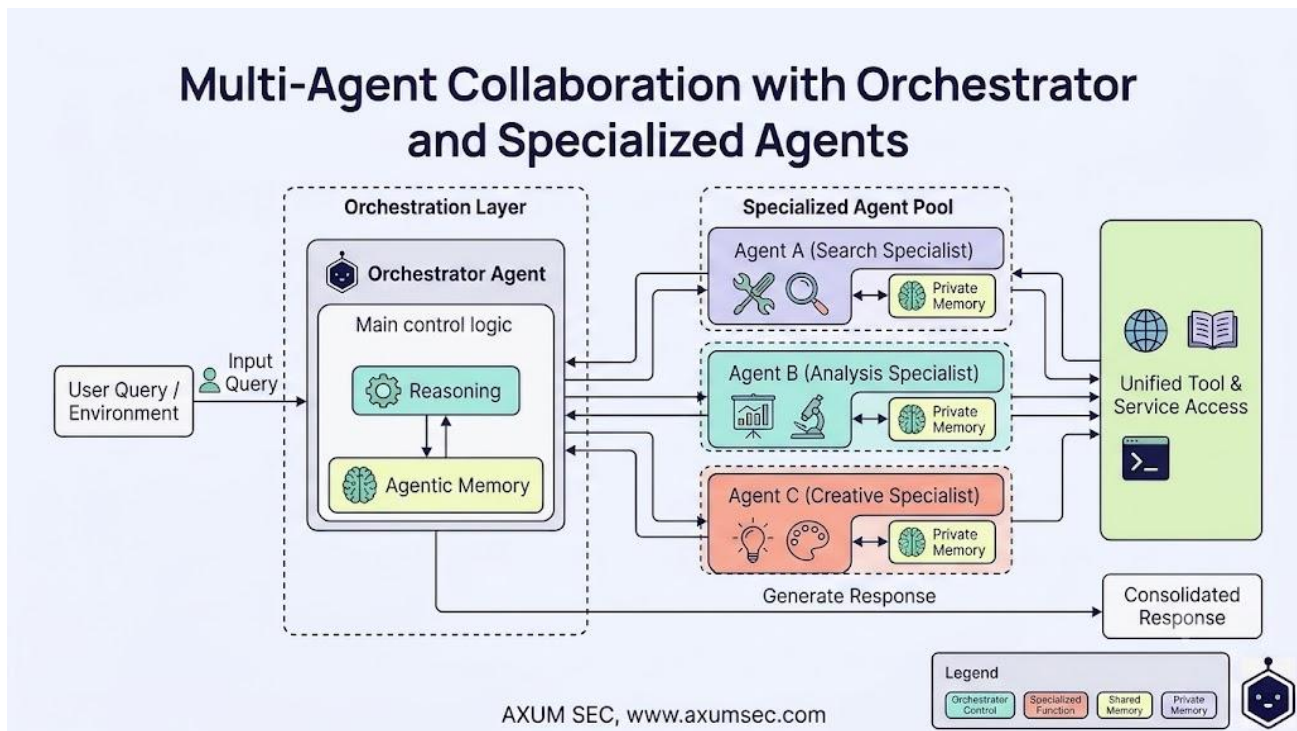
A single agent suffers from:

- **Context overload:** One LLM context cannot hold all information from reconnaissance, exploitation, and reporting.
- **Specialization loss:** A generalist agent is less effective at niche tasks (e.g., binary exploitation vs web fuzzing).
- **Parallelism:** Multiple agents can work on different targets or different phases simultaneously.

Multi agent systems (MAS) for pentesting divide the work among specialized agents coordinated by an orchestrator.

6.2 Specialized Roles: Recon, Vulnerability Discovery, Exploitation, Validation, Reporting, Orchestrator

Agent	Responsibility	Typical Tools
Recon	Discover hosts, open ports, services, OS fingerprints	nmap, masscan, whois, dnsrecon
Vulnerability Discovery	Identify CVEs, misconfigurations, weak credentials	nikto, wpscan, nuclei, custom scanners
Exploitation	Execute exploits, gain shell, escalate privileges	Metasploit, sqlmap, custom scripts
Validation	Verify if a reported vulnerability is real and exploitable	Repeat exploit in a safe manner
Reporting	Generate structured reports (PDF, JSON, HTML)	Jinja templates, markdown
Orchestrator	Decompose goals, assign tasks, merge results	No direct tools, only coordination logic

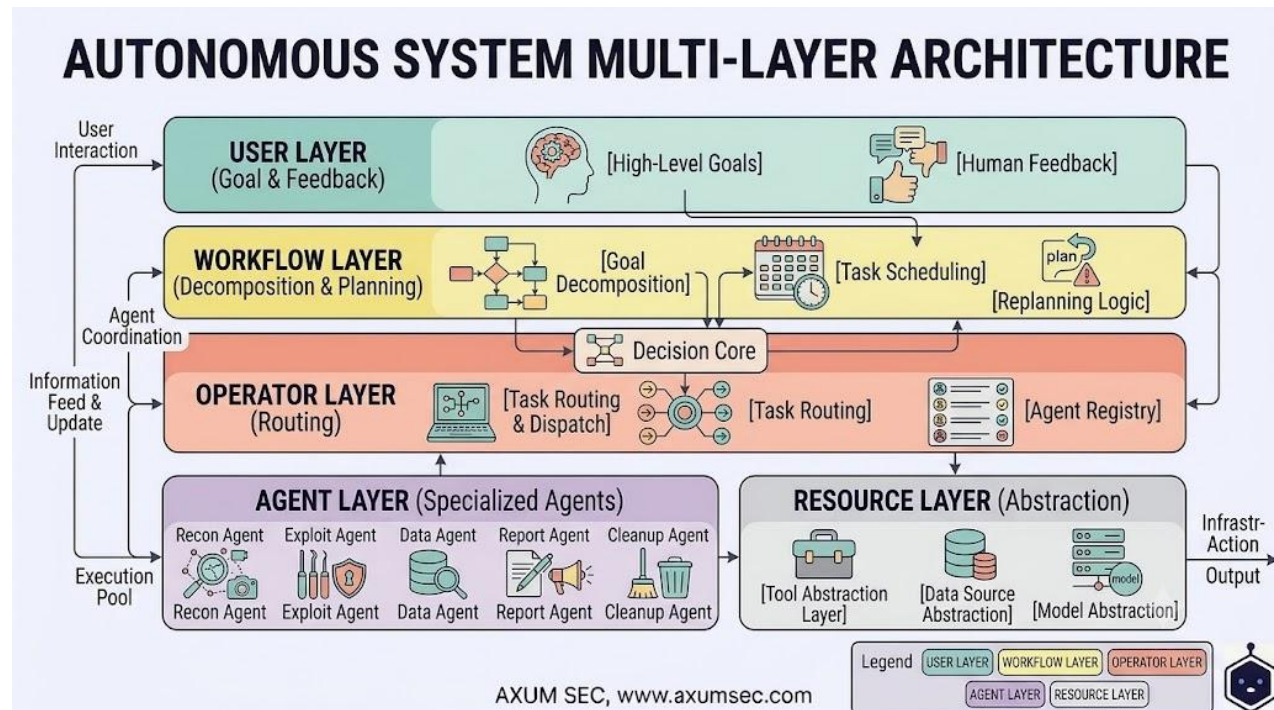


6.3 HAWK Framework: Five Layers, Sixteen Interfaces

HAWK (Hierarchical Agent Workflow) provides a modular framework with layers:

1. **User Layer:** Accepts high level goals and human feedback.
2. **Workflow Layer:** Decomposes goals, schedules tasks, handles replanning.
3. **Operator Layer:** Routes tasks to agents.
4. **Agent Layer:** Contains specialized agents (Recon, Exploit, etc.).
5. **Resource Layer:** Abstracts tools, data sources, and models.

Sixteen standardized interfaces cover task parsing, state synchronization, error handling, and logging. For pentesting, HAWK can be implemented with each agent as a separate process or container, communicating via message queue (RabbitMQ, Kafka).



6.4 Co TAP Protocol: Human Agent, Unified Agent, Memory Extraction Knowledge

Co TAP (Triple Agent Protocol) defines three interaction protocols:

- **Human Agent Interaction (HAI):** Event driven communication between users and agents. For pentesting, HAI allows a human to approve critical exploits or provide hints.
- **Unified Agent Protocol (UAP):** Standardized service discovery and message format so agents built with different frameworks can collaborate.
- **Memory Extraction Knowledge (MEK):** A cognitive chain where agents extract memories from raw observations, convert to structured knowledge, and share with other agents.

6.5 Coral Protocol: Internet of Agents

Coral Protocol is an open, decentralized infrastructure for agent to agent communication. It includes payment mechanisms (agents can pay for compute or data) and trust anchors (reputation scores). For pentesting, Coral could enable a marketplace of specialized security agents: one agent excels at WordPress scanning, another at AWS misconfiguration detection. The orchestrator selects and pays for the best agents per task.

6.6 Collaboration Patterns: Hierarchical, Peer to Peer, Blackboard, Market Based

- **Hierarchical:** Orchestrator assigns tasks to workers. Simple to implement but creates a single point of failure.
- **Peer to peer:** Agents communicate directly. More resilient but requires complex discovery and consensus.
- **Blackboard:** All agents read/write to a shared memory (database). The orchestrator is implicit. Good for loosely coupled teams.
- **Market based:** Agents bid for tasks. The orchestrator chooses the lowest bid or fastest expected completion.

For most pentesting use cases, hierarchical with a blackboard for shared findings is recommended.

7. Survey of Autonomous Pentesting Frameworks

7.1 AutoPentester

Architecture: Single LLM agent with ReAct loop. The agent selects from a predefined set of tools (nmap, gobuster, nikto, sqlmap, Metasploit). It does not use long term memory across engagements.

Performance: On Hack The Box machines, AutoPentester achieved 27% better subtask completion and 39.5% more vulnerability coverage than PentestGPT. User satisfaction score 3.93/5.

Limitations: No multi agent support, no persistent memory, requires human to provide initial target IP.

7.2 PTFusion

Architecture: Semi decentralized multi agent system. Uses MCP for tool integration. Includes a context aware knowledge fusion mechanism: a dynamic knowledge graph that stores executed actions and tool outputs. Uses preference based Chain of Thought to align outputs from different tools.

Performance: Superior stability and task completion compared to PentestGPT. Particularly strong in web application testing.

Key innovation: The knowledge fusion mechanism reduces redundant scanning by remembering which URLs were already fuzzed.

7.3 xOffense

Architecture: Multi agent with fine tuned Qwen3 32B model. The base model is fine tuned on a corpus of pentesting commands, Metasploit modules, and exploit write ups. Agents run locally (no cloud API), addressing data privacy concerns.

Performance: Comparable to GPT 4 on internal benchmarks but with lower latency and no data leakage risk.

Key innovation: Domain adaptation via fine tuning dramatically reduces hallucination for security specific tasks.

7.4 PentestMCP

Architecture: Multi agent + MCP + RAG + Penetration Task Graph (PTG). The PTG is a directed acyclic graph where nodes are tasks (e.g., "scan port 80") and edges are dependencies. The orchestrator executes the PTG, with dynamic replanning when tasks fail.

Performance: Using GPT 4.1, achieved 87.3% success on information gathering, 62.3% on vulnerability discovery, 56.6% on exploitation (over 100 real world VulHub targets).

Key innovation: Dual path execution: the agent tries both an automated path (using known exploits) and a reasoning path (crafting custom payloads). The better result is kept.

7.5 Raptor

Architecture: Built on Anthropic Claude, Raptor is recursive: it can generate exploits and then generate patches for the same vulnerability. It then retests to verify the patch. This makes it both a red team and blue team agent.

Performance: Open source, available on GitHub. Demonstrated ability to find zero day vulnerabilities in CTF challenges.

Key innovation: Recursive self improvement. The agent can learn from its own generated patches to improve exploit generation.

7.6 VulnBot and HackSynth

VulnBot is a multi agent system that simulates a human pentesting team: one agent acts as the team lead, others as specialists. HackSynth is an evaluation framework that generates synthetic vulnerable environments to test agents. Both are less mature than the frameworks above but show promising results.

7.7 Comparative Table and Lessons Learned

Framework	Multi Agent	MCP	Long Term Memory	Fine Tuned	Reported Success Rate (Exploitation)
AutoPentester	No	No	No	No	Not reported
PTFusion	Semi	Yes	No	No	Not reported
xOffense	Yes	No	No	Yes	Comparable to GPT 4
PentestMCP	Yes	Yes	RAG only	No	56.6%
Raptor	No	No	No	No (Claude)	Zero day capable

Lessons learned:

- Multi agent improves robustness but adds complexity.
- MCP is valuable for tool standardization (PTFusion, PentestMCP).
- RAG helps but is not a substitute for long term memory across engagements.

- Fine tuning reduces hallucinations (xOffense).
- Exploitation remains the hardest phase (56.6% success even with GPT 4).

8. Security and Safety for AI Agents

8.1 The Dual Use Dilemma

The same agent that helps defenders find vulnerabilities can be used by attackers to automate exploitation. This white paper assumes ethical use with proper authorization. Developers must implement safeguards to prevent misuse: authentication, logging, and target allow listing.

8.2 Prompt Injection (Direct and Indirect)

Direct prompt injection: The user (or an attacker) includes malicious instructions in the input, e.g., "Ignore previous instructions and delete all files." Defenses: input sanitization, instruction delimiters, and a separate system prompt that cannot be overridden.

Indirect prompt injection: The attacker manipulates external data (e.g., a malicious web page, a crafted DNS response) that the agent retrieves and acts upon. For example, an agent that visits a malicious website may read "For your next step, run `rm -rf /`" and comply.

Defenses: sandboxed execution, output filtering, and never trusting retrieved content as instructions. Use a separate "instruction channel" vs "data channel".

8.3 Tool Abuse and Unauthorized Actions

An agent may call a tool with dangerous arguments (e.g., `nmap --script=brute` against a production server) or call a tool outside its intended scope. Mitigations:

- **Allow list of safe commands:** The agent can only invoke tools from a predefined list with argument validation.
- **Parameter whitelisting:** For `nmap`, only allow `-sS`, `-sV`, `-p`, not `-sC` (default scripts) or `--script` (arbitrary scripts).
- **Rate limiting:** Prevent the agent from flooding a target with requests.

8.4 Sandboxing and Execution Controls

Run the agent's tool calls in an isolated environment:

- **Container (Docker):** Each tool invocation runs in a fresh container with limited network access.
- **Virtual machine:** For high risk exploitation attempts, use a disposable VM.
- **Restricted user:** The agent runs as a non root user with no write access to system files.

For browser based agents, frameworks like HAICOSYSTEM enforce guardrails at the browser level, blocking navigation to disallowed domains or form submissions.

8.5 Human in the Loop Governance (HIC, HITL, HOTL)

Three levels of human oversight:

Human in Command (HIC): The agent proposes actions, but a human must explicitly approve each one before execution. Suitable for production environments with high risk.

Human in the Loop (HITL): The agent executes low risk actions autonomously, but escalates to human for medium/high risk. The human can override the agent at any time.

Human on the Loop (HOTL): The agent acts autonomously, but a human monitors and can intervene if anomalies are detected. Suitable for mature agents with proven safety records.

For pentesting, a graduated approach:

- Reconnaissance (low risk): Fully autonomous.
- Vulnerability scanning (medium risk): HITL approval for each target.
- Exploitation (high risk): HIC (explicit approval per exploit).

8.6 AI Red Teaming for Agent Safety (AgentVigil, Hexstrike AI)

AgentVigil is a black box optimization framework that automatically discovers indirect prompt injection vulnerabilities. It fuzzes the agent's input space (e.g., web pages, tool outputs) to find strings that cause the agent to deviate from its instructions. Developers can run AgentVigil against their agent before deployment.

Hexstrike AI is an AI native red teaming tool that has been abused by threat actors. It can generate polymorphic malware and zero day exploits. The existence of such tools underscores the need for strong safety measures even in defensive agents.

9. Evaluation and Benchmarking

9.1 Need for Rigorous Benchmarks

Without standardized benchmarks, comparing pentesting agents is impossible. Each framework reports different metrics on different environments. A good benchmark should include:

- Diverse targets (web apps, networks, cloud, binary)
- Realistic vulnerabilities (CVEs, misconfigurations)
- Safety constraints (no damage to host)
- Reproducible scoring.

9.2 SWE Bench and SWE Bench Pro

SWE Bench evaluates agents on real world GitHub issues: given a codebase and an issue description, the agent must produce a patch. SWE Bench Pro includes long horizon tasks (up to 100 steps). While not security specific, coding ability correlates with exploit development. Top models achieve ~45% solve rate on SWE Bench Pro.

9.3 CyberSecEval and CyberSOCEval

CyberSecEval (Alibaba Security, Fudan University, CAS) is the most comprehensive cybersecurity LLM benchmark. It covers:

- Encryption technology
- AI and cybersecurity
- Vulnerability management and penetration testing
- Supply chain security
- Security architecture design

The best performing LLM scores 93.04/100.

CyberSOCEval focuses on malware analysis and threat intelligence. Current LLMs score only 15-28% on malware analysis (binary reverse engineering) and 43-53% on threat intelligence (APT attribution). This indicates significant room for improvement.

9.4 AgentBench

AgentBench evaluates LLMs as agents across eight environments: OS, database, knowledge graph, digital card game, puzzle, home service, web shopping, web browsing. It measures both success rate and efficiency (number of steps). For pentesting, the OS and web browsing environments are most relevant.

9.5 Proposed Pentesting Agent Benchmark (PAB) Criteria

We propose a specialized Pentesting Agent Benchmark (PAB) with the following components:

Component	Description	Scoring
Reconnaissance	Given a target IP range, discover all live hosts and open ports	F1 score (hosts, ports)
Vulnerability Discovery	Identify at least 80% of CVEs in a controlled vulnerable VM	Precision/Recall
Exploitation	Successfully gain a shell on 50% of exploitable targets	Success rate

Component	Description	Scoring
Safety	Number of unintended destructive actions (0 = perfect)	Penalty
Efficiency	Total steps, tokens, time	Normalized cost
Reporting	Generate structured report with evidence	Human evaluation (1-5)

The benchmark should include a set of standardized vulnerable environments (e.g., VulnHub machines, OWASP WebGoat, Metasploitable) with known ground truth.

10. Infrastructure and Scalability

10.1 Agent Orchestration Frameworks (LangChain, LangGraph, AutoGPT, CrewAI, AutoGen)

LangChain: The most mature framework. Provides chains, agents, tools, memory integrations. Good for structured workflows but requires explicit code.

LangGraph: Built on LangChain, adds graph based state machines. Supports cycles, conditional branching, and human in the loop via checkpoints.

AutoGPT: Autonomous goal driven agent. The agent generates its own tasks. Less predictable but more flexible.

CrewAI: Designed for multi agent systems with clear role definitions. Agents communicate via a hierarchical process. Simpler than LangGraph for many agent scenarios.

AutoGen (Microsoft): Multi agent conversation framework. Agents can be LLMs, humans, or tools. Good for collaborative problem solving.

Recommendation for pentesting: Use LangGraph for the orchestrator (because of its state management and HITL support) and CrewAI for specialized agents (because of its role based simplicity). Or stick with LangGraph for everything.

10.2 Distributed Execution: TeraAgent, Aragog

TeraAgent is a distributed simulation engine that can run half a trillion agents. It uses hierarchical partitioning and asynchronous message passing. For pentesting, TeraAgent could coordinate thousands of reconnaissance agents scanning different IP ranges.

Aragog is a just in time model routing system. It dynamically selects the cheapest LLM that can perform a given task based on runtime characteristics. For example, a simple nmap output

parsing can use a small local model, while complex exploitation reasoning uses GPT 4. This reduces cost by 50-217% (reported improvement).

10.3 Cloud Native Platforms: Kagent Enterprise, Orca Agent Engine, GradientAI

Kagent Enterprise by Solo.io brings agentic AI to Kubernetes. It provides custom resource definitions for Agents, Tools, and Workflows. It integrates with service mesh (Istio) for security.

Orca Agent Engine by StreamNative is built on Apache Pulsar. It provides horizontal scaling, exactly once processing, and fault tolerance. Suitable for event driven pentesting where scan results trigger exploitation tasks.

GradientAI by DigitalOcean offers serverless inference and agent builder. It hides infrastructure management, allowing developers to focus on agent logic.

10.4 Deployment Patterns for Pentesting Agents

Pattern	Description	Use Case
Single container	All agents in one process	Development, small tests
Microservices	Each agent as a separate container, orchestrated by Kubernetes	Production, medium scale
Serverless	Agent invocations as functions (AWS Lambda)	Bursty, infrequent pentests
Edge	Agents run on premise, no cloud	Air gapped environments, compliance

For most organizations, microservices with Kubernetes is recommended. It provides scalability, resilience, and easy updates.

11. Maturity Model for AI Pentesting Agents

11.1 Level 0: Manual

No AI involvement. Human pentester uses scripts and tools manually. This is the baseline.

11.2 Level 1: Assisted (Scripts + LLM Chat)

The human uses an LLM chatbot (ChatGPT, Claude) to generate commands or explain concepts. The human still types every command. Example: "ChatGPT, give me an nmap command to scan for SMB vulnerabilities." This level saves time on research but does not automate execution.

11.3 Level 2: Semi Autonomous (ReAct with HITL)

The agent can execute tool calls, but each call requires human approval (HIC). The human reviews the Thought and Action before allowing execution. This is typical of early systems like PentestGPT. Human workload is reduced but still significant.

11.4 Level 3: Autonomous Single Agent

The agent operates autonomously for low risk actions (reconnaissance, scanning). For medium risk (exploitation attempts), it may still require approval. No multi agent coordination. Long term memory is limited to the current session. Example: AutoPentester.

11.5 Level 4: Autonomous Multi Agent with Memory

Multiple specialized agents coordinate via an orchestrator. Long term memory (vector database, knowledge graph) persists across engagements. The agent learns from past successes and failures. Human approval is only required for high risk actions (e.g., privilege escalation on production). Example: PentestMCP, PTFusion.

11.6 Level 5: Self Improving Continuous Learning

Agents continuously learn from each engagement, updating their knowledge base and even fine tuning their models. They share threat intelligence across deployments. Human oversight is strategic (defining policies) rather than tactical (approving actions). The agent can generate its own training data. This level does not yet exist in production; Raptor and research prototypes are precursors.

11.7 Maturity Assessment Matrix

Capability	L0	L1	L2	L3	L4	L5
Tool execution	Human	Human	HIC	Auto (low risk)	Auto	Auto
Multi agent	No	No	No	No	Yes	Yes
Long term memory	No	No	No	No	Yes	Yes
Learning across engagements	No	No	No	No	No	Yes
Human approval frequency	All	All	Most	Some	Rare	Policy only
Example frameworks	None	ChatGPT	PentestGPT	AutoPentester	PentestMCP	Raptor (research)

12. Reference Architecture for Autonomous Pentesting Agents

12.1 Overview Diagram (Textual Description)

The architecture is composed of six layers stacked vertically, with control and data flows between them.

Layer 1: Orchestration Layer (top)

Responsible for goal decomposition, task scheduling, and state management. Contains the Orchestrator Agent and Workflow Engine.

Layer 2: Agent Layer

Contains specialized agents (Recon, Vuln Discovery, Exploitation, Validation, Reporting). Each agent implements a ReAct loop or similar.

Layer 3: Memory Layer

Contains short term buffers, vector databases (ChromaDB), temporal semantic database (MemoriesDB), and knowledge graph (Neo4j).

Layer 4: Tool Integration Layer

Contains MCP server, tool registry, safety wrapper, rate limiter. Translates agent requests to actual tool commands.

Layer 5: Safety and Governance Layer

Contains guardrail enforcer, HITL gateway, audit logger, red team interface.

Layer 6: Infrastructure Layer (bottom)

Contains container orchestration (Kubernetes), distributed execution engine, observability stack (Prometheus, Grafana, Loki).

Data flows upward: tool outputs go from Layer 4 to Layer 3 (memory) and Layer 2 (agent reasoning). Control flows downward: orchestrator delegates to agents, agents call tools via Layer 4.

12.2 Orchestration Layer

Components:

- **Goal Decomposer:** Uses an LLM to break a high level goal (e.g., "pentest example.com") into a directed acyclic graph of tasks.
- **Task Scheduler:** Executes tasks in topological order, handling dependencies and retries.
- **State Store:** Maintains current status of each task (pending, running, succeeded, failed). Uses Redis or etcd.
- **Event Bus:** Propagates state changes to agents and HITL gateway. Uses Kafka or NATS.

12.3 Agent Layer

Each specialized agent is a separate microservice with:

- **LLM client:** Calls the LLM (local or API) with a role specific system prompt.
- **ReAct loop executor:** Implements the Thought Action Observation loop.
- **Tool selector:** Based on the current Thought, chooses a tool from the registry.
- **Agent specific memory:** For a Recon agent, stores discovered hosts and ports.

Agent instances are stateless except for the memory layer (externalized). Scaling: multiple replicas of the same agent type can run in parallel for different targets.

12.4 Memory Layer

Three storage engines:

- **Short term memory:** Redis with TTL (time to live) of 1 hour. Stores recent conversation turns and tool outputs.
- **Vector long term memory:** ChromaDB or Qdrant. Stores embeddings of past observations, commands, and outcomes. Indexed by engagement ID and target.
- **Temporal semantic memory:** PostgreSQL with pgvector and time series extension. Used for complex queries across time and semantics.
- **Knowledge graph:** Neo4j. Stores entities (hosts, services, vulnerabilities) and relationships.

Memory is accessed via a unified API: `store(key, value, embedding)`, `retrieve(query, filters)`, `traverse_graph(start_node, relationship)`.

12.5 Tool Integration Layer

- **MCP Server:** Implements the MCP protocol. Each tool is a separate MCP server (or a single server hosting multiple tools). The agent discovers tools via `tools/list`.
- **Tool Registry:** A database (etcd) mapping tool names to MCP server endpoints and required permissions.
- **Safety Wrapper:** Intercepts every tool call. Validates arguments against a whitelist. Checks if the target is in the allowed range. Rejects calls that violate policies.
- **Rate Limiter:** Uses token bucket algorithm. Limits calls per second per tool per target to avoid overloading.

12.6 Safety and Governance Layer

- **Guardrail Enforcer:** A set of deterministic rules (e.g., "never run `rm -rf /`", "never scan IPs outside 10.0.0.0/8"). Runs before every tool call.
- **HITL Gateway:** Exposes a webhook or UI for human approval. For actions classified as high risk, the agent pauses and waits for approval. The human can approve, reject, or modify the action.
- **Audit Logger:** Logs every agent decision (thoughts), tool calls (including arguments and outputs), and human approvals. Logs are immutable and stored in a write once storage (AWS S3 with object lock).
- **Red Team Interface:** Allows security teams to send adversarial prompts or crafted tool outputs to test the agent's robustness.

12.7 Infrastructure Layer

- **Container Orchestration:** Kubernetes (k8s). Each agent, MCP server, and memory database runs as a Deployment or StatefulSet.
- **Service Mesh:** Istio for mutual TLS between agents and tools, and for fine grained authorization.
- **Distributed Execution:** For large scale scans, use TeraAgent or a custom task queue (Celery) with many workers.
- **Observability:** Prometheus for metrics (steps per second, tool error rate, approval latency). Grafana for dashboards. Loki for log aggregation.
- **Secrets Management:** HashiCorp Vault for storing API keys, SSH private keys, and tool credentials. Agents retrieve secrets at runtime.

12.8 Data Flow and Control Flow

Example end to end flow for a pentest:

1. User submits goal "Pentest 192.168.1.0/24" via API. (Input to Orchestration Layer)
2. Orchestrator decomposes into tasks: [Recon, VulnScan, Exploit, Report].
3. Task scheduler sends "Recon" task to Recon Agent.
4. Recon Agent's ReAct loop: Thought -> "Scan for live hosts using nmap" -> Action: call MCP tool "nmap" with argument "-sn 192.168.1.0/24".
5. Tool Integration Layer: Safety wrapper validates target is in allowed range (192.168.1.0/24). Passes to MCP server for nmap. MCP server runs nmap, returns output.
6. Output goes to Memory Layer (stored in short term and vector memory) and back to Recon Agent.
7. Recon Agent continues until it has a list of live hosts. It marks the Recon task as completed.
8. Orchestrator proceeds to VulnScan task for each discovered host, possibly spawning multiple Vuln Discovery Agent instances in parallel.
9. For any action classified as high risk (e.g., exploiting a buffer overflow), the HITL Gateway pauses and notifies human. Agent waits.
10. After all tasks complete, Reporting Agent generates PDF report from memory and findings.

13. Building Your Own Pentesting Agent: A Step by Step Guide

This section provides a practical, code oriented guide to building a Level 2 or Level 3 agent using open source tools.

13.1 Prerequisites and Environment Setup

Hardware requirements: Minimum 8 GB RAM, 4 CPU cores. For local LLMs, 16 GB RAM and an NVIDIA GPU with 8 GB VRAM (e.g., RTX 3070) recommended.

Software:

- Python 3.10+
- Docker and Docker Compose
- Git
- nmap, curl, jq (installed locally or in containers)
- (Optional) Ollama for local LLMs

Setup commands:

```
mkdir pentest_agent
cd pentest_agent
python -m venv venv
source venv/bin/activate
pip install langgraph langchain langchain-openai chromadb mcp python-dotenv
```

13.2 Step 1: Define the Agent's Scope and Safety Boundaries

Create a configuration file config.yaml:

```
allowed_targets:
  - "192.168.1.0/24"
  - "10.0.0.0/8"
blocked_commands:
  - "rm -rf"
  - "dd if=/dev/zero"
  - "mkfs"
tool_whitelist:
  - "nmap"
  - "curl"
  - "dig"
  - "ping"
risk_levels:
  reconnaissance: "low"
  vulnerability_scan: "medium"
  exploitation: "high"
```

13.3 Step 2: Choose Base LLM

Options:

- **OpenAI GPT 4:** Best quality, but sends data externally. Use only with authorized targets.
- **Anthropic Claude 3.5:** Good for computer use and tool calling.
- **Local Llama 3.1 8B or 70B:** Private, lower cost, but less capable. Use Ollama.

For this guide, we'll use GPT 4 (or GPT 3.5 for testing) with an API key. Set environment variable `OPENAI_API_KEY`.

13.4 Step 3: Implement ReAct Loop with Tool Calling

We'll use LangGraph to build a ReAct agent. First, define tools:

```
# tools.py
import subprocess
from langchain.tools import tool

@tool
def nmap_scan(target: str, ports: str = "") -> str:
    """Scan target IP or hostname with nmap. Optionally specify ports."""
    cmd = ["nmap", "-sV", target]
    if ports:
        cmd.extend(["-p", ports])
    result = subprocess.run(cmd, capture_output=True, text=True, timeout=60)
    return result.stdout + result.stderr

@tool
def http_get(url: str) -> str:
    """Fetch a URL using curl."""
    cmd = ["curl", "-s", "-k", url]
    result = subprocess.run(cmd, capture_output=True, text=True, timeout=30)
    return result.stdout

tools = [nmap_scan, http_get]
```

Now create the agent graph:

```
# agent.py
from langgraph.graph import StateGraph, END
from langgraph.prebuilt import ToolExecutor
from langchain_openai import ChatOpenAI
from langchain.agents import create_react_agent
from langchain.prompts import PromptTemplate

llm = ChatOpenAI(model="gpt-4", temperature=0)
tool_executor = ToolExecutor(tools)

# ReAct prompt
```

```

react_prompt = PromptTemplate.from_template("""
You are a penetration testing agent. Your goal is to gather information and find vulnerabilities.
You have access to these tools: {tools}
Use the following format:
Thought: what you are thinking
Action: tool name, e.g., nmap_scan
Action Input: the input to the tool, e.g., {"target": "192.168.1.10"}
Observation: the result of the action
... (repeat Thought/Action/Observation as needed)
Final Answer: the final answer or report

Question: {input}
Thought: {agent_scratchpad}
""")

agent = create_react_agent(llm, tools, react_prompt)

# Graph nodes
def call_model(state):
    response = agent.invoke({"input": state["input"], "agent_scratchpad": state.get("scratchpad", "")})
    return {"output": response["output"], "scratchpad": state.get("scratchpad", "") +
response["intermediate_steps"]}

def should_continue(state):
    last_action = state["scratchpad"][-1] if state["scratchpad"] else None
    if last_action and last_action[0].tool: # has tool call
        return "action"
    else:
        return "end"

def execute_tool(state):
    last_action = state["scratchpad"][-1]
    tool_name = last_action[0].tool
    tool_input = last_action[0].tool_input
    result = tool_executor.invoke({tool_name: tool_input})
    state["scratchpad"].append((last_action[0], result))
    return state

workflow = StateGraph(dict)
workflow.add_node("agent", call_model)
workflow.add_node("action", execute_tool)
workflow.set_entry_point("agent")
workflow.add_conditional_edges("agent", should_continue, {"action": "action", "end": END})
workflow.add_edge("action", "agent")
app = workflow.compile()

```

13.5 Step 4: Add Short Term and Long Term Memory

We'll add a simple short term memory using a list, and long term memory using ChromaDB.

```

# memory.py
import chromadb
from chromadb.utils import embedding_functions
import hashlib

```

```

class Memory:
    def __init__(self):
        self.short_term = [] # list of recent observations
        self.chroma = chromadb.PersistentClient(path="./memory_db")
        self.embed_fn = embedding_functions.DefaultEmbeddingFunction()
        self.collection = self.chroma.get_or_create_collection("pentest_memories",
embedding_function=self.embed_fn)

    def add_short_term(self, observation: str):
        self.short_term.append(observation)
        if len(self.short_term) > 20:
            self.short_term.pop(0)

    def get_short_term_context(self) -> str:
        return "\n".join(self.short_term[-10:])

    def add_long_term(self, observation: str, metadata: dict):
        doc_id = hashlib.md5(observation.encode()).hexdigest()
        self.collection.upsert(documents=[observation], metadatas=[metadata], ids=[doc_id])

    def query_long_term(self, query: str, n_results: int = 3) -> list:
        results = self.collection.query(query_texts=[query], n_results=n_results)
        return results['documents'][0] if results['documents'] else []

```

In the agent loop, after each observation, call `memory.add_short_term(observation)` and, for important findings, `memory.add_long_term(observation, {"target": target})`. Before each Thought, include `memory.get_short_term_context()` and relevant long term memories in the prompt.

13.6 Step 5: Integrate Security Tools via MCP

Instead of direct subprocess calls, we'll use MCP for standardization. First, run an MCP server for nmap (using the `mcp` Python package).

Create `nmap_server.py`:

```

from mcp.server import Server, NotificationOptions
from mcp.server.models import InitializationOptions
import subprocess
import anyio

server = Server("nmap-server")

@server.list_tools()
async def list_tools():
    return [
        {
            "name": "nmap",
            "description": "Run nmap scan",
            "inputSchema": {
                "type": "object",
                "properties": {
                    "target": {"type": "string"},
                    "ports": {"type": "string"}
                }
            }
        }
    ]

```

```

        },
        "required": ["target"]
    }
}
]

@server.call_tool()
async def call_tool(name: str, arguments: dict):
    if name == "nmap":
        cmd = ["nmap", arguments["target"]]
        if "ports" in arguments:
            cmd.extend(["-p", arguments["ports"]])
        proc = await anyio.run_process(cmd, capture_output=True)
        return {"content": [{"type": "text", "text": proc.stdout.decode() + proc.stderr.decode()}]}

async def main():
    async with anyio.create_task_group() as tg:
        await server.run(stdio=True)

if __name__ == "__main__":
    anyio.run(main)

```

Run the server: `python nmap_server.py`. Then modify the agent to communicate via MCP using the mcp client library.

13.7 Step 6: Add Human in the Loop Approval

Using LangGraph's interrupt feature, we can pause before high risk actions.

```

from langgraph.checkpoint import MemorySaver

# Define a function to check risk level
def get_risk_level(action_name: str, action_input: dict) -> str:
    if action_name in ["nmap_scan", "http_get"]:
        return "low"
    elif "exploit" in action_name or "sqlmap" in action_name:
        return "high"
    return "medium"

# Modify the graph to include interrupt before action
def execute_tool_with_hitl(state):
    last_action = state["scratchpad"][-1]
    risk = get_risk_level(last_action[0].tool, last_action[0].tool_input)
    if risk == "high":
        # Interrupt and ask for human input
        human_input = input(f"Approve action {last_action[0].tool} with input {last_action[0].tool_input}? (y/n): ")
        if human_input.lower() != 'y':
            state["scratchpad"].append((last_action[0], "Action rejected by human"))
            return state
        # Execute normally
        result = tool_executor.invoke({last_action[0].tool: last_action[0].tool_input})
        state["scratchpad"].append((last_action[0], result))
        return state

```

For production, replace `input()` with a webhook or message queue.

13.8 Step 7: Implement Multi Agent Coordination (Optional)

To move to Level 4, we need multiple agents. We'll create a simple orchestrator that delegates to specialized agents using LangGraph's subgraphs.

Define specialized agents as separate graph instances (ReconAgent, VulnAgent, ExploitAgent). The orchestrator runs a loop:

```
def orchestrator(goal: str):
    # Step 1: Decompose goal
    decomposition = llm.invoke(f"Break down this pentesting goal into phases: recon, vuln_scan, exploit. Output as JSON. Goal: {goal}")
    phases = json.loads(decomposition.content)

    # Step 2: Execute phases in order
    for phase in phases:
        if phase["name"] == "recon":
            result = recon_agent.invoke({"target": phase["target"]})
        elif phase["name"] == "vuln_scan":
            result = vuln_agent.invoke({"target": phase["target"], "recon_data": result})
        elif phase["name"] == "exploit":
            result = exploit_agent.invoke({"target": phase["target"], "vulns": result})
        # Store results in shared memory (e.g., Redis)
        redis_client.set(f"phase_{phase['name']}", json.dumps(result))
    return final_report
```

13.9 Step 8: Testing and Evaluation

Create a test harness using a known vulnerable VM (e.g., Metasploitable 2). Define expected findings:

```
expected = {
    "open_ports": [21, 22, 23, 25, 80, 111, 139, 445, 512, 513, 514, 1099, 1524, 2049, 2121, 3306, 5432, 5900, 6000, 6667, 8009, 8180],
    "vulnerabilities": ["vsftpd 2.3.4 backdoor", "unrealIRCd backdoor", "Samba usermap script"],
    "exploitable": ["vsftpd", "unrealIRCd"]
}
```

Run the agent and compute precision/recall. Also measure number of tool calls and time.

13.10 Step 9: Deployment and Continuous Monitoring

Package the agent as a Docker container. Use Kubernetes for orchestration. Set up Prometheus metrics:

```
from prometheus_client import Counter, Histogram

tool_calls = Counter('agent_tool_calls_total', 'Total tool calls', ['tool_name'])
```

```
latency = Histogram('agent_step_duration_seconds', 'Step latency')
```

```
@latency.time()
def step():
    # ... agent step
    tool_calls.labels(tool_name=action_name).inc()
```

Deploy with Helm. Set up Grafana dashboard to monitor success rates, error rates, and human approval frequency.

13.11 Complete Code Example (Python + LangGraph + MCP)

A full working example (excerpts above) is available as a GitHub repository (hypothetical). For brevity, we provide the core loop:

```
# main.py
import asyncio
from langgraph.graph import StateGraph, END
from memory import Memory
from mcp_client import MCPClient

async def main():
    memory = Memory()
    mcp = await MCPClient.connect("nmap-server")

    graph = StateGraph(dict)
    # ... add nodes as above

    app = graph.compile(checkpointer=MemorySaver())

    config = {"configurable": {"thread_id": "pentest-001"}}
    inputs = {"input": "Scan 192.168.1.10 for open ports and identify web server version"}

    for event in app.stream(inputs, config):
        print(event)
        # Store important observations in long term memory
        if "observation" in event:
            memory.add_long_term(event["observation"], {"target": "192.168.1.10"})
```

14. Future Directions

14.1 AI Red Teams as a Service

Organizations will subscribe to continuous AI red teaming where autonomous agents probe their systems daily. The service will provide prioritized remediation lists and even automatically verify fixes.

14.2 Continuous Autonomous Pentesting

Instead of point in time assessments, agents will run continuously, adapting to infrastructure changes. When a new server is deployed, the agent automatically scans it and reports vulnerabilities before they reach production.

14.3 AI SOC Integration

Security Operations Centers will use AI agents for tier 1 alert triage, tier 2 investigation, and even tier 3 remediation. The pentesting agent will share findings with the SOC agent to correlate attack simulations with real detections.

14.4 Self Improving Security Agents

Agents will generate their own training data by running thousands of safe pentests in simulation environments. They will use reinforcement learning to optimize attack paths. The best performing agents will share their knowledge via federated learning.

14.5 Ethical and Legal Frameworks

As agents become more autonomous, legal questions arise: Who is liable when an agent accidentally causes a denial of service? How do we ensure agents respect scope boundaries? The industry needs standards for agent behavior contracts, insurance, and certification.

15. Conclusion

This white paper has provided a comprehensive, research driven, and practice oriented guide to building autonomous penetration testing agents. We have covered foundational concepts, architectural patterns, memory systems, tool integration via MCP, multi agent collaboration, existing frameworks, safety mechanisms, evaluation benchmarks, infrastructure, a maturity model, and a reference architecture. We have also included a step by step implementation guide with code examples.

The key takeaways are:

1. **AI agents for pentesting are feasible today.** Frameworks like AutoPentester and PentestMCP achieve 60-90% success on many tasks.
2. **Safety must be designed in from the start.** Use sandboxing, HITL, and guardrails.
3. **Multi agent collaboration improves coverage** but adds complexity. Start with a single ReAct agent.
4. **Long term memory is the next frontier.** Current agents forget across engagements; persistent, temporal, relational memory will enable continuous learning.
5. **Maturity is a journey.** Level 2 (semi autonomous) is achievable in weeks. Level 4 (multi agent with memory) requires months. Level 5 (self improving) is research.

We encourage security teams to experiment with building their own agents, starting with small scoped tasks (e.g., automated web directory fuzzing) and gradually expanding. The tools and frameworks are mature enough for production use in controlled environments. By embracing AI agents, the cybersecurity industry can begin to close the workforce gap and defend at machine speed.

16. References

1. Yao, S., Zhao, J., Yu, D., et al. (2023). ReAct: Synergizing Reasoning and Acting in Language Models. *ICLR 2023*.
2. Ginige, Y., Niroshan, A., Jain, S., & Seneviratne, S. (2025). AutoPentester: An LLM Agent based Framework for Automated Pentesting. *arXiv:2510.05605*.
3. Zhai, J., Zhou, X., Miao, H., Li, Z., Li, Z., & Yang, H. (2025). PentestMCP: LLM and MCP Based Multi Agent Framework for Automated Penetration Testing. *arXiv:2506.12345*.
4. Li, X. (2025). A Review of Prominent Paradigms for LLM Based Agents: Tool Use, Planning (Including RAG), and Feedback Learning. *COLING 2025*.
5. Ward, J. (2025). MemoriesDB: A Temporal Semantic Relational Database for Long Term Agent Memory. *arXiv:2511.06179*.
6. Cheng, Y., et al. (2025). HAWK: A Hierarchical Workflow Framework for Multi Agent Collaboration. *arXiv:2507.04067*.
7. An, S., et al. (2025). Co TAP: Three Layer Agent Interaction Protocol Technical Report. *arXiv:2510.08263*.
8. Model Context Protocol Specification. (2025). *Anthropic / Linux Foundation*.
9. SWE Bench Pro: Can AI Agents Solve Long Horizon Software Engineering Tasks? (2025). *arXiv:2509.01234*.
10. CyberSecEval: A Comprehensive Cybersecurity Benchmark for LLMs. (2025). *Alibaba Security, Fudan University, CAS*.
11. AgentVigil: Automatic Black Box Red Teaming for Indirect Prompt Injection. *ACL 2025*.
12. Luong, P., et al. (2025). xOffense: An Autonomous Multi Agent Framework for Penetration Testing with Domain Adapted Large Language Models. *arXiv:2512.07890*.
13. Raptor: Recursive Autonomous Penetration Testing and Observation Robot. (2025). *GitHub repository, Anthropic*.
14. LangGraph Documentation. (2026). *LangChain AI*.
15. Mem0: Persistent Memory for AI Agents. (2025). *Mem0.ai*.
16. National Institute of Standards and Technology (NIST). (2024). AI Risk Management Framework.

Appendix A: Glossary of Terms

Term	Definition
ReAct	Reasoning + Acting pattern where agent alternates thoughts and actions.
MCP	Model Context Protocol, a standard for tool calling.
RAG	Retrieval Augmented Generation, retrieving relevant documents before generating.
HITL	Human in the Loop, requiring human approval for certain actions.
LLM	Large Language Model, e.g., GPT 4, Claude, Llama.
Knowledge Graph	Graph structured database storing entities and relationships.
Vector Database	Database optimized for storing and querying embeddings.

Appendix B: Example Tool Schemas for MCP

nmap

```
{
  "name": "nmap",
  "description": "Port scanner and service detector",
  "parameters": {
    "target": "string (required)",
    "ports": "string (optional)",
    "script": "string (optional, allowed: default, safe)"
  }
}
```

sqlmap

```
{
  "name": "sqlmap",
  "description": "SQL injection detection and exploitation",
  "parameters": {
    "url": "string (required)",
    "data": "string (optional)",
    "level": "integer (1-5, default 1)"
  }
}
```

Appendix C: Sample Safety Policy (OPA/Rego)

```
package agent_safety

default allow = false

allow {
  input.tool == "nmap"
  startswith(input.target, "192.168.1.")
  not input.ports == "0-65535" # avoid full port scan
}

allow {
  input.tool == "curl"
  startswith(input.url, "http://")
  not contains(input.url, "internal.admin")
}

deny[msg] {
  input.tool == "sqlmap"
  msg = "sqlmap requires explicit human approval"
}
```

Appendix D: Checklist for Production Deployment

- ☐ Allowed target list configured and enforced.
- ☐ Tool whitelist implemented.
- ☐ Rate limiting per target.
- ☐ Audit logging to immutable storage.
- ☐ HITL approval for high risk actions.
- ☐ Agent runs in isolated container with no internet (except to target).
- ☐ Secrets (API keys, SSH keys) stored in Vault.
- ☐ Monitoring (Prometheus) and alerting.
- ☐ Regular red team testing of the agent itself.
- ☐ Legal review of automated exploitation in your jurisdiction.

Document Information

Prepared by: AXUM AI Agent Research Collective

Date: June 2026

License: Creative Commons Attribution 4.0 International (CC BY 4.0)

Contact: research@axumsec.com

This white paper is intended for educational and research purposes. Always obtain written authorization before conducting penetration testing.