

Agentic-AI Powered Continuous Offensive Security

RidgeGen™: Autonomous Security Validation at Machine Speed



Agentic AI Powered Red Teaming & Pentesting

Multi-Step Attack Chains

Zero-Hallucination Verification

Three-Zone Safety Guardrails

Model-Agnostic Engine

Business-Logic Flaw Discovery

Click-Away Remediation

On-Premises & Data-Private

RidgeGen™ brings red-team depth to every security team

RidgeGen™ is an agentic AI-powered platform that augments security teams by autonomously executing exploitation to continuously validate your organization's posture. Where traditional automated penetration testing tools stop at detection, RidgeGen™ reasons through multi-step attack chains, surfaces business-logic flaws, and uncovers unknown risks at machine speed, and fixes them with a click-away remediation plan.

Built on a proprietary harness framework layered over AI models, RidgeGen™ pairs a multi-layered memory design and a multi-model architecture with a unique autonomous verification process. The result is explainable, repeatable, and deterministic findings—each one validated with exploitable evidence and paired with specific, actionable remediation.

RidgeGen™ shortens the mean time from exposure to validated, exploitable detection and to remediation from months to minutes, all while operating within strict, enforceable safety guardrails that protect production environments.

Challenges

Modern attack surfaces grow faster than security teams can test them. Vulnerability scanners overwhelm security teams with unprioritized findings, while legacy automated pentest tools often deliver limited results that create a false sense of security by missing exploitable attack paths. Meanwhile, skilled red-team talent remains scarce and expensive. Business-logic flaws and chained, multi-step attack paths—precisely the issues that lead to real breaches—routinely slip past script or catalog-based tools.

At the same time, security leaders are wary of pointing autonomous, offensive tooling at live systems. You need depth and autonomy without losing control over what an agent is allowed to touch, and you need defensible evidence rather than the unpredictable output that general-purpose AI models produce on their own.

From exposure to validated exploit—and to remediation— in minutes, not months.

RidgeGen™ Solution & Key Benefits

RidgeGen™ turns the non-deterministic nature of AI models into a definitive, defensible path. It maintains test paths, plans, scope, evidence, and validation throughout an engagement, producing reasoning with human-like maturity while keeping every action auditable. An open prompt interface invites your domain knowledge and human creativity directly into the agent, learning your specific environment and capturing institutional expertise in a shared enterprise memory.

Key Benefits

- Bring red-team capability in-house and reduce reliance on costly third-party engagements
- Validate which vulnerabilities are actually exploitable in your unique environment
- Uncover business-logic and unknown risks that traditional tools miss
- Cut mean time to remediation from months to minutes
- Continuously test web applications before release

Platform Differentiators

- 1. Model-agnostic flexibility:** A model-agnostic design lets teams choose and switch between supported frontier or self-hosted models backends to balance cost, performance, and data-residency requirements.
- 2. Zero hallucination:** A unique autonomous verification process confirms every finding with reproducible, exploitable evidence before it is reported.
- 3. Greater efficiency:** The harness framework drives results using fewer tokens than general-purpose model usage, lowering operating cost.
- 4. More findings:** A proprietary multi-layered memory design enables deeper, longer reasoning chains and surfaces issues other approaches overlook.
- 5. Harness-based guardrails:** Safety and accuracy are enforced by the harness—outside model discretion—not left to the model to self-police.

How RidgeGen™ Works

RidgeGen™ is driven entirely through natural-language prompts. Users describe the target, scope, and objective in plain language—either as a single paragraph or through a guided conversation—and the agent plans, executes, verifies, and explains the engagement from end to end.

Prompt-Based Configuration

A conversational interface gathers everything needed to define an assessment: the target URL or IP address, whether to include subdomains or assets in subnets, which pages or paths to include or exclude, login credentials and MFA for authenticated testing, the testing objective (speed vs. depth), specific risk classes of interest, and the purpose of the test—such as PCI external testing, OWASP compliance, third-party or supplier penetration testing, or continuous testing. The agent stays focused on testing scope and confirms configuration before any run begins.

Prompt-Based Result Interpretation

A floating AI chat window is available at any point during or after a test. Users can ask questions in natural language and receive grounded answers drawn from trustworthy resources—including the CISA catalog, the NVD, Exploit Database, the OWASP project, and compliance standards such as SOC 2, PCI DSS, HIPAA, CMMC, and ISO 27001—plus the engagement's own logs, timestamps, and findings.

Users can compare results across time and against past runs, request side-by-side comparisons and trend analysis, and ask practical, role-aware questions such as where to start remediation, how to fix a specific issue, or which findings to present to leadership.

Autonomous Exploitation & Verification

RidgeGen™ reasons through multi-step attack chains, escalates where evidence supports it, and validates each finding through its autonomous verification process. Findings are presented with the supporting evidence needed to reproduce them, eliminating false positives and the manual triage they create.

Specific, Actionable Remediation

RidgeGen™ builds remediation guidance on top of code-level evidence rather than generic advice—delivering click-away remediation that points teams directly to the fix and supports rapid patch verification.

Security Guardrails — Safety by Design

Because RidgeGen™ is an autonomous offensive agent with real reach into customer systems, its guardrails are treated as a product safety contract. Every safety-relevant boundary is governed by trusted controls, enforced outside model discretion, and recorded as proof for customers and inspectors. RidgeGen™ demonstrates real weaknesses with exploitable evidence while avoiding disruption or unintended impact to your environment.

Three Operating Zones			
Zone	Green	Amber	Red
Posture	Safe	Supervised	Authorized-Aggressive
Permitted Actions	No write or delete actions are performed against targets. Recommended for sensitive production systems.	No delete or modification of existing data. A temporary write may occur where required to prove a finding.	Write, modify, and delete actions may occur and may be unrecoverable. Recommended for lab and pre-production testing.

Safebox — Your Credentials Stay Yours

All user-supplied credentials and sensitive information—usernames, passwords, and PII—are held in a protected Safebox.

- Safebox is enforced in every zone.
- Safebox contents are never passed to any third-party system, including model vendors.
- Sensitive values can be masked from the GUI and from generated reports.
- Credentials are used in memory only during a test and are not persisted in RidgeGen™.
- Any information written to disk is encrypted, and evidence is limited to high-level database structure or table metadata—never customer database contents.

Model-Agnostic by Design

RidgeGen™'s multi-model architecture lets users select and switch between supported model backends from a single configuration page. Supported backends span frontier models—including Google Gemini, Anthropic Claude, and OpenAI—self-hosted open-source models that keep testing entirely within your environment, as well as AI model gateways and aggregators such as OpenRouter, LiteLLM, and Amazon Bedrock. This model-agnostic approach means teams can optimize for cost, capability, or data residency, and customers may also bring their own model into the platform. Selecting and switching models is an administrator-level control.

Deployment

On-premises deployment keeps testing, findings, and sensitive data fully within your environment—ideal for organizations with strict data-residency and confidentiality requirements.

RidgeGen™ System Requirements

RidgeGen™ deploys on-premises as a self-contained software package. The following are the minimum requirements for the on-premises deployment.

On-Premises Hardware Requirements

Component	Specification
Processor (CPU)	8-core CPU
Memory	32 GB RAM
Storage	200 GB
Base System	Ubuntu 24.04.4 LTS
Concurrent Tasks	Up to 4 concurrent testing tasks at the minimum configuration (excluding LLM-side processing)

Supported AI Model Providers

RidgeGen™ is model-agnostic and connects to leading model providers through their APIs or an OpenAI-compatible endpoint. Because specific model versions are released and updated continuously, individual models are not listed here; the current selection is always available on the AI Models configuration page.

- Anthropic Claude
- Google Gemini
- Google Vertex AI
- OpenAI (GPT)
- DeepSeek
- OpenAI-compatible endpoints — for self-hosted and open-source models, keeping testing within your environment
- AI Model Gateways & Aggregators — such as OpenRouter, LiteLLM, and Amazon Bedrock for unified model routing and management

RidgeGen™ Detailed Features

Interface & Experience

- Modern, simplified agentic-AI interface aligned with Ridge Security branding
- Prompt-based configuration and natural-language interaction throughout
- Visualized kill-chains and attack maps for clear understanding of the attack path
- Real-time progress indicators—progress bars and target beacons—so users always see the engagement advancing
- Configuration history tab that retains the key points captured from the chat for later retrieval

Scenario Templates

Built-in, one-click scenarios accelerate common engagements, including

- Authenticated Web Testing
- API Testing
- Host-Based Testing
- Windows Active Directory Testing

Reporting & Compliance

- Reports enriched with the test plan and attack-surface topology
- Real-time multi-language support across the GUI, events, findings, and downloadable reports
- Compliance-oriented report options including PCI DSS, CIS, NIST, and HIPAA

Framework Mapping & Discovery

- Automatic mapping of findings to the MITRE ATT&CK™ framework with modern visual presentation
- Automatic mapping of vulnerabilities to OWASP Top 10 categories across web and API targets
- Misconfiguration included as a first-class testing objective
- Passive and active asset and network-topology discovery, capturing device type, firmware version, and topology to enrich offensive context

Operations & Management

- User-selectable AI model backend—choose and switch between Google Gemini, Anthropic Claude, OpenAI, and self-hosted open-source models to match cost, capability, and data-residency needs
- Scheduled and recurring tasks defined in natural language (for example, weekly or twice-weekly cadences)
- Token management with per-task usage, aggregate cost monitoring, and configurable limits
- Two-tier user management with super-admin and regular-user roles and backend data separation per user
- Online and offline upgrades, including skip-version upgrades, with source-code encryption protecting platform integrity
- Externally published API documentation for third-party integration

Licensing

RidgeGen™ offers flexible licensing for both enterprise end-users and managed security service providers (MSSPs). The agent always confirms license usage with the user before a run begins.

Licensing Type

Units of Consumption

Web Testing License	One license per unique FQDN tested
API Testing License	One license per API test (URL, API endpoint, or uploaded Swagger file)
Host Testing License	One license per unique IP address tested

Models

- **End-user annual subscription:** Each license binds to a FQDN or IP address and allows unlimited testing of that FQDN for the license term (typically one year from activation).
- **MSSP quota license:** An annual pool of licenses; each run deducts licenses based on the number of IPs, FQDNs, or API files tested.

RidgeGen™ — depth, validated, and safe by design
Autonomous exploitation at machine speed, with evidence you can defend.

Visit Our Website



Follow Us On LinkedIn



About Ridge Security Technology Inc.

Ridge Security empowers CISOs to build cyber resilience through AI-native continuous offensive security. By helping organizations continuously validate their security posture, focus on proven cyber risks, and accelerate remediation, Ridge Security enables security teams to stay ahead of evolving threats while maximizing the effectiveness of their existing security investments. Recognized by Gartner and Frost & Sullivan, Ridge Security serves enterprises worldwide across financial services, government, telecommunications, and other critical industries.



Ridge Security Technology Inc.

<https://ridgesecurity.ai/>



© 2026 All Rights Reserved Ridge Security Technology Inc. RidgeBot® and RidgeGen™ are registered trademarks of Ridge Security Technology Inc.