

AI PROVING GROUNDS

Test & Validate **AI Agents**

Establish trust in your AI-empowered SOC. Continuously stress-test your human and agentic defenses against LLM threats fueled by Mythos and Daybreak.

How do security leaders prove that newly deployed AI agents will behave as expected when a Mythos- or Daybreak-fueled zero-day hits? Trust cannot be based on assumptions. Organizations need a risk-free AI agent testing infrastructure to continuously interrogate agent and vendor performance, confirm policy boundaries, and measure operational safety.

The Trust Blueprint: Rigorous Agentic AI Testing & Validation

Safely launch cyberattacks directly against autonomous agent architectures to assess their decision-making logic, map failure thresholds, and establish trust in your AI-empowered SOC.

1. Install the Agent in a Realistic Range

Rapidly host prospective third-party or internal AI agents inside a zero-risk, digital replica range.

2. Conduct Adversarial AI Agent Stress-testing

Execute automated, continuous attack loops to evaluate response logic, prioritization accuracy, and both malware & living-off-the-land (LoTL) detection.

3. Score Performance

Deliver unified event scoring, impact maps, and run-based reporting to assess the readiness of your agents.

Key Outcomes



Objective Readiness Signals

Discover and quantify if agents perform exactly as anticipated or introduce unintended structural vulnerabilities.



A Custom Framework for Trust

Define, quantify, and track security boundaries tailored specifically to your organization's unique operational tolerances.

Build Trust in Your AI.

Talk to a SimSpace AI security expert. [SimSpace.com](https://www.SimSpace.com)