

AI Proving Grounds **Playbook**

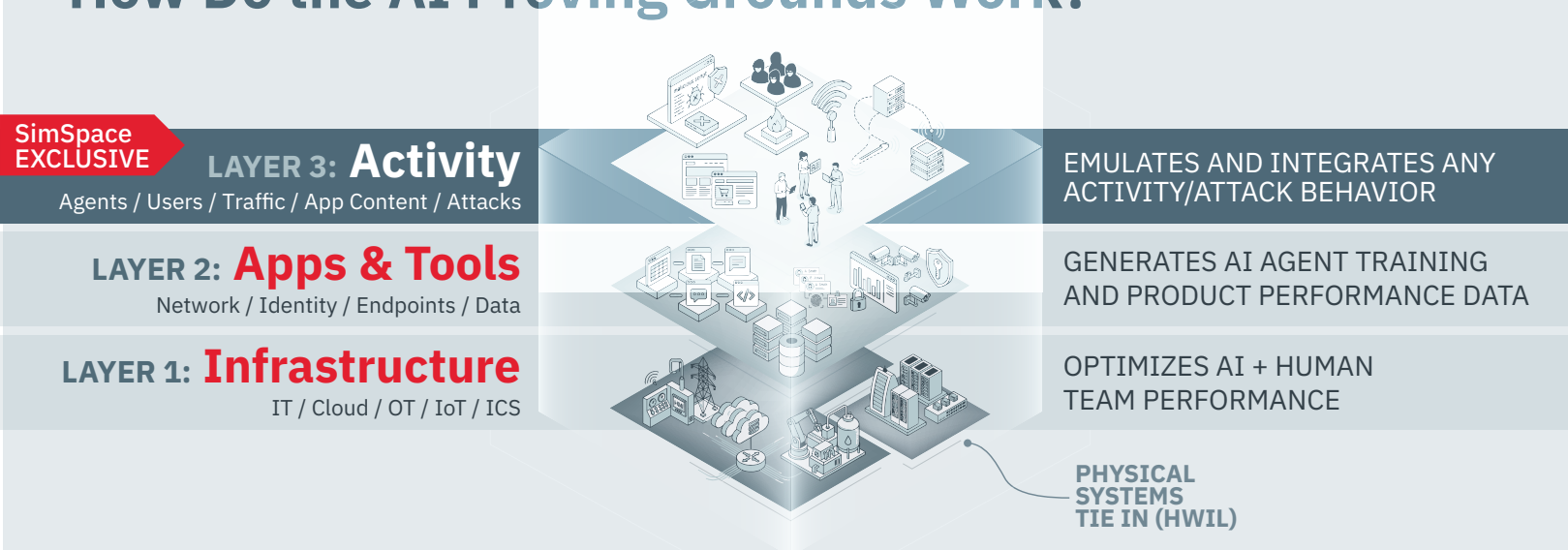
Train, Test, and Operationalize AI Agents to Defeat AI-Fueled Adversaries

AI agents rarely fail in clean lab environments. They fail in real security operations with noisy SOC workflows and adversary pressure.

SimSpace is the AI Proving Grounds where you can continuously train and test your AI agents against LLM threats fueled by modern AI models like Mythos and Daybreak.

Build trust that your AI agents will hold up under real-world pressure.

How Do the AI Proving Grounds Work?



Generate Hyper-Synthetic data

Generate hyper-synthetic telemetry and attack data based on true enterprise conditions to make models and agents smarter and more predictive. Then, retrain with production data, all without impacting production environments.

Test & Validate Agentic Solutions

Test and validate autonomous AI agents operating within full workflows. Validate agentic performance results against realistic user and attack emulation and adversary behavior inside a replica of your environment.

Prove Agentic Solutions Are Safe & Trustworthy

Evaluate an agent's policy adherence, escalation thresholds, and failure modes in a realistic cyber simulation infrastructure to build operational trust and ensure AI decisions are explainable and secure.

Strengthen Human-Agent Coordination

By training and testing inside the AI Proving Grounds, organizations can properly evaluate human and AI integration to prove that both AI agents and human operators can stand up against real threats.

Build & Train AI Agents with Hyper-Synthetic Data

Stop training AI agents on sterile, theoretical datasets.
Generate hyper-synthetic data to build & train red and blue agents that are ready for real-world attacks.

When product teams and corporate data science units train AI agents on static data, they inherit fragile automation logic, model drift, and a high rate of false positives. Autonomous combat-ready AI agents need live, precision-labeled, full-stack log data generated in a modern AI training software ecosystem.

SimSpace allows us to create dynamic datasets over and over again, different types of attacks, changing the environment to really replicate our different customers, so we can train our agents without risking customer data.

– Senior Product Manager, Top 3 AI Hyperscaler

Build Real-World-Ready Agents with Hyper-Synthetic Data for AI Training

The SimSpace AI Proving Grounds is an isolated, realistic cyber simulation platform that generates hyper-synthetic, production-grade log data to build resilient, context-aware red and blue agents.

1. Execute Targeted Adversary Emulation

Launch nation-state and e-crime attacks within an isolated, enterprise-mirrored environment to force genuine tool interactions and telemetry.

2. Harvest Multi-Vector Log Data

Automatically export comprehensive log data to data science and product management teams across both security stack architectures and everyday productivity app traffic.

3. Refine Agent Performance in Cyber Simulations

Provide data scientists with the real-world, hyper-synthetic log metrics to train agent model candidates against realistic background conditions.

4. Benchmark & Validate AI Agents

Place newly minted agents back into the live-fire scenarios to map out and quantify iterative step-by-step performance developments.

Key Outcomes



Compressed Time-to-Market

Eradicate weeks spent hand-crafting or cleaning target data sets.



Bulletproof AI Validation & Product Verification

Deploy security features with proven metrics to achieve real-world operational effectiveness.



Proven Agents

Move smoothly from alpha builds to battle-hardened agents capable of identifying complex threat signatures.

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.

Coordinate human-AI SecOps playbooks and AI-driven detection engineering.

SecOps excellence requires a hybrid architecture: human intuition acting alongside tactical AI machine speed. Yet, integrating AI into production workflows without practicing team coordination risks catastrophic alert confusion, missed handoffs, and tool friction. Enterprises must operationalize AI agents and human analysts as a single, unified unit that rehearses—and defends—together.

Enhance Human-Agent Collaboration in Your AI-Empowered SOC

Operationalize combined human-AI SecOps playbooks within a digital replica of your actual production infrastructure with full-stack processes and real data and workflows.

- 1. Install the Agent in a Realistic Range**

Rapidly host AI agents inside a zero-risk, digital replica range where human operators can coordinate alongside the agents.
- 2. Conduct Adversarial AI Agent Stress-testing**

Execute automated, continuous attack loops to evaluate response logic, prioritization accuracy, and lateral movement detection.
- 3. Attack the Human-Agent Team**

Deliberately attempt prompt injections, data tampering, and model-exploitation techniques to map out unique security vulnerabilities before they hit your production environment.
- 4. Score Performance**

Deliver unified event scoring, impact maps, and run-based reporting to assess the readiness of your agents.
- 5. Optimize Across Playbook Scenarios**

Score operational handoffs across complex incident response scenarios, refining prompt instructions and escalation handshakes until they are flawless.

Key Outcomes

- Maximized ROI on Modern SecOps**

Fine-tune automation playbooks to slash alert triage timelines and minimize false positive bottlenecks.
- Verified Human-AI Collaboration & Synergy**

Eradicate operational gaps, confirm that team-to-agent escalations occur exactly as planned, and confidently deploy to production.

Build Trust in Your AI.

Whether your team is building responsible AI agents, validating purchased AI tools, or operationalizing AI agents alongside your human operators, the AI Proving Grounds helps your team build trust in its agentic SOC.

Ready to put the AI Proving Grounds playbook into action for your security operations?

Talk to an AI security expert at SimSpace. [SimSpace.com](https://simspace.com)