

ThreatSpace™: Defend the AI Pipeline

Benefits:

- Practice hands on hunting for and investigating AI-specific attacks in a realistic, application environment.
- Learn to analyze raw AI application telemetry and correlate them with Google Security Operations (SecOps) to track an attack lifecycle.
- Clarify shared responsibility roles with AI by understanding the distinction between a foundational models' security and your organizational responsibilities to protect custom agent workflows.
- Analyze recent threat intelligence, real-world case studies, and verified research Proof of Concepts (PoCs) from the past 12 months to validate your cyber defense strategies with today's threat landscape.

This hands-on workshop covers the core security principles, attack vectors, and proactive safeguards necessary to secure internally developed, AI-enabled applications and agentic workflows. Rather than focusing on abstract model training, security teams will learn how to defend the entire application integration surface—from user-facing prompts to downstream vector databases, APIs, and autonomous system actions.

Learners are invited into the state-of-the-art ThreatSpace™ cyber range for the practical experience of uncovering security incidents and analyzing forensic artifacts targeting AI application infrastructure. Throughout the course, learners will gain hands-on experience investigating threat actor activity targeting orchestration layers, Retrieval-Augmented Generation (RAG) sources, and high-privilege agents, allowing them to apply these defensive skills directly in their daily operations.

ThreatSpace™ is an engaging cyber range featuring a virtualized enterprise within Google Cloud. In this delivery, security professionals access a virtualized cloud range hosting a custom, internally developed AI application. This hands-on environment lets learners work directly with the application's runtime logs, RAG data sources, and system prompts, correlating these artifacts within Google Security Operations (SecOps) to investigate real-world application-layer attacks.

The workshop covers the complete **OWASP Top 10 for LLM Applications**, addressing the full spectrum of risks identified in real-world enterprise deployments. By investigating these vulnerabilities in a live environment, teams learn how to trace attacker behaviors, analyze application-layer logs, and audit their overall security posture against modern, AI-driven threats.

OWASP Top 10 for LLM Applications (2025)

- | | |
|---|--|
| <ul style="list-style-type: none"> • LLM01: Prompt Injection • LLM02: Sensitive Information Disclosure • LLM03: Supply Chain Vulnerabilities • LLM04: Data and Model Poisoning • LLM05: Improper Output Handling | <ul style="list-style-type: none"> • LLM06: Excessive Agency • LLM07: System Prompt Leakage • LLM08: Vector and Embedding Weaknesses • LLM09: Misinformation • LLM10: Unbounded Consumption |
|---|--|