

AI Penetration Testing

Secure Your AI Before It's Exploited

Offensive AI Techniques We Simulate

- Prompt injection and manipulation
- Agent exploitation and impersonation
- Data poisoning and model extraction
- Workflow hijacking and tool abuse

What You Get

- Executive overview and risk summary
- Detailed technical findings for remediation teams
- AI-specific threat modeling
- Proof-of-concept attack scenarios
- Prioritized remediation roadmap to quickly reduce risk

Standalone or Integrated

This offering can be delivered as a standalone AI Penetration Test or a milestone within our larger Security Assessment Suite

Get Ahead of AI Threats

Be proactive, not reactive. Schedule your AI Penetration Test and secure your AI before it is exploited.

**To learn more about our AI services,
contact your Account Team today!**

1.800.998.0067

www.connection.com/services

AI adoption is accelerating faster than most security programs were built to support. While traditional penetration testing focuses on applications, networks, and infrastructure, AI introduces new attack vectors that bypass conventional controls.

Our AI Penetration Test is a purpose-built offensive security engagement designed to uncover real-world vulnerabilities in your AI ecosystem before attackers do.

Why This Matters

AI attacks are evolving faster than traditional defenses can adapt. Organizations must now defend against:

- Prompt injection attacks that manipulate model behavior and outputs
- Compromised AI decision-making through agent and workflow abuse
- Data leakage through embeddings and retrieval systems
- Hidden backdoors across the AI supply chain

Without targeted testing, these attack paths often go undetected. This assessment ensures your AI systems are secure, resilient, and trusted.

What We Test

We simulate how modern attackers target AI-enabled environments, including:

- Generative AI and LLM applications
- AI agents and autonomous workflows
- Retrieval-Augmented Generation (RAG) pipelines and knowledge systems
- Machine learning models and embeddings
- AI infrastructure across cloud platforms, APIs, and containers
- AI supply chain, including datasets, models, and integrations

Built on Advanced AI Red Teaming Expertise

Our methodology is aligned with AI offensive security practices, including:

- AI system reconnaissance and mapping
- Multi-agent exploitation techniques
- RAG and embedding attacks
- AI infrastructure compromise scenarios

How We Test

Using an AI-focused red team methodology, we:

- Map your AI attack surface
- Identify high-value AI assets and trust boundaries
- Exploit vulnerabilities across AI systems
- Simulate real-world attack paths
- Deliver actionable remediation guidance

