

EC-Council
Building A Culture Of Security

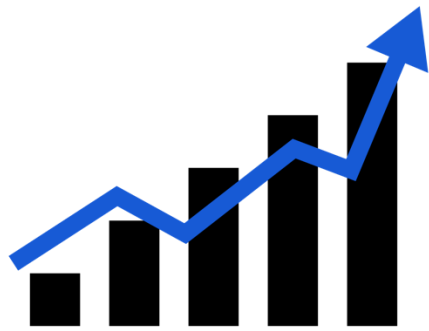
C|OASP
Certified Offensive AI Security Professional

*Certified Offensive
AI Security Professional*

Emulate. Exploit. Defend.

The New *Reality*

Attackers are using AI to create and weaponize threats faster than security teams can respond



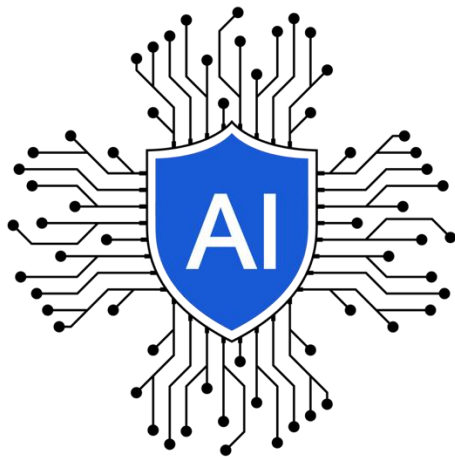
270%

increase in AI-powered
cyberattacks since 2020,
Hoplon Infosec Report,
2024



\$10.5_T

projected cybercrime cost by
2025, Cybersecurity Ventures
Report



77%

of organizations lack the
foundational data and
AI security practices,
Accenture, 2025



USD
670_K

added breach cost due to
shadow AI. IBM Report, 2025

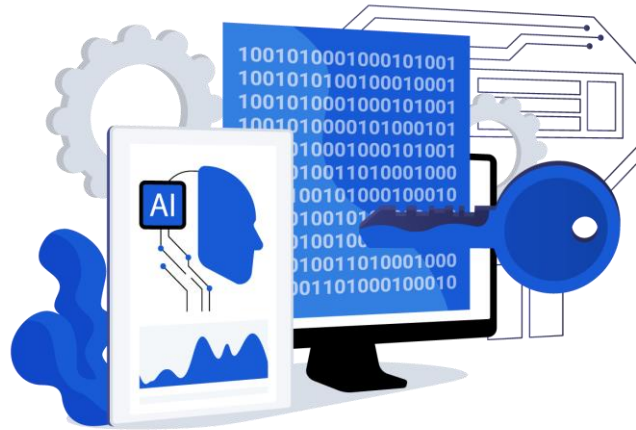
The New *Reality*

Attackers are using AI to create and weaponize threats faster than security teams can respond



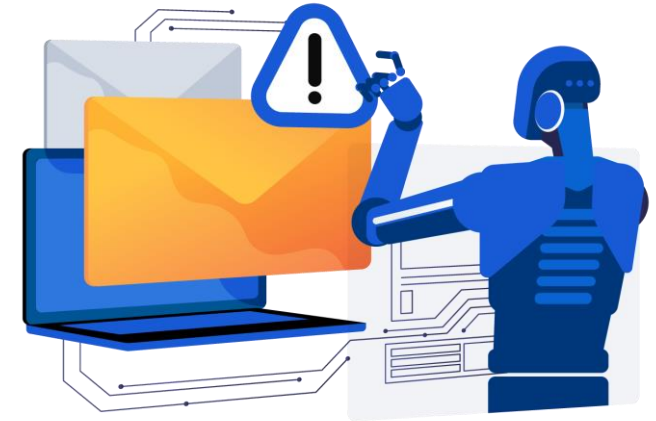
87%

of organizations worldwide
report being hit by AI-powered
cyberattacks in the past year



78%

of CISOs say that AI-powered
threats are having a significant
impact on their organizations



Only
35%

of leaders feel they have
prepared employees
effectively for AI roles.

Market-wide *Challenges*



○ **AI systems: From Chatbots to autonomous Agents** - Now power business operations, but adversarial attacks are outpacing defenses.

○ Malicious prompt injections, context manipulation, and model theft are **triggering data leaks and outages**.

○ Compromised supply chains and unsafe plugins expose **new attack surfaces**

○ Most organizations **lack experts** who can think like attackers within AI systems or validate them through structured risk assessments.

○ Traditional cybersecurity teams aren't equipped to test LLMs, agent architectures, or data pipelines for adversarial threats, **leaving costly AI projects stalled by risk and compliance gaps**.

AI Security Is Now a *C-Suite Priority*

AI security spend is expanding

Global end-user spending on information security is expected to reach **\$213B in 2025** (up from **\$193B in 2024**) and hit **\$240B in 2026**.

Board/C-suite forcing action

By 2026, **64%** of companies assess AI risks before deployment (up from **37%** previous year).

AI spending boom = larger attack surface

Worldwide AI spending forecast to total **\$2.52 trillion in 2026 (+44% YoY)**.

AI Red Teaming is emerging as a category

AI Red Teaming Services market valued at **\$1.75B (2025)** → projected **\$4.8B (2029) (28.6% CAGR)**.

AI cybersecurity is becoming its own market

AI cybersecurity solutions market: **\$30.92B (2025)** → projected **\$86.34B (2030) (22.8% CAGR)**.

Red teaming still rare = high demand opportunity

Less than **20%** have adopted practices like **regular red teaming exercises** (despite AI governance growth).

The Gap: Skilled Owners to *Drive Secure AI Programs*



AI security skill gap inside cyber teams

41% of cybersecurity leaders cite **AI as the most pressing skills need** for their teams (higher than cloud security).



Lack of AI security assessment capability

66% expect AI to have the most significant impact on cybersecurity, but only **37%** have processes to assess AI tool security before deployment.



Skill shortages block AI security adoption

54% cite **skill shortages** as a key challenge in adopting AI for cybersecurity.



AI cybersecurity talent shortage

34% of cybersecurity professionals say their companies lack workers with **AI cybersecurity capabilities**.

A New Generation of *AI Defenders is Emerging Across Industries*

However, securing AI requires red-team-to-blue-team capability: to understand models, exploit weaknesses, then build safeguards.

Because adversarial threats affect every AI deployment, the industries where this role is most impactful include:



Finance & Banking

For fraud detection models, trading algorithms, and customer-data models vulnerable to theft or poisoning.



Healthcare

AI diagnostics, imaging, and patient data privacy, where model leakage or manipulation can cause serious real-world harm.



Government, Defense, & Critical Infrastructure

SCADA, autonomous systems, and national security data where attackers may exploit AI agents or plugins.



Tech/Cloud Platforms/AI Providers

Providers of AI platforms, agents, or plugins need to defend their products and customers.

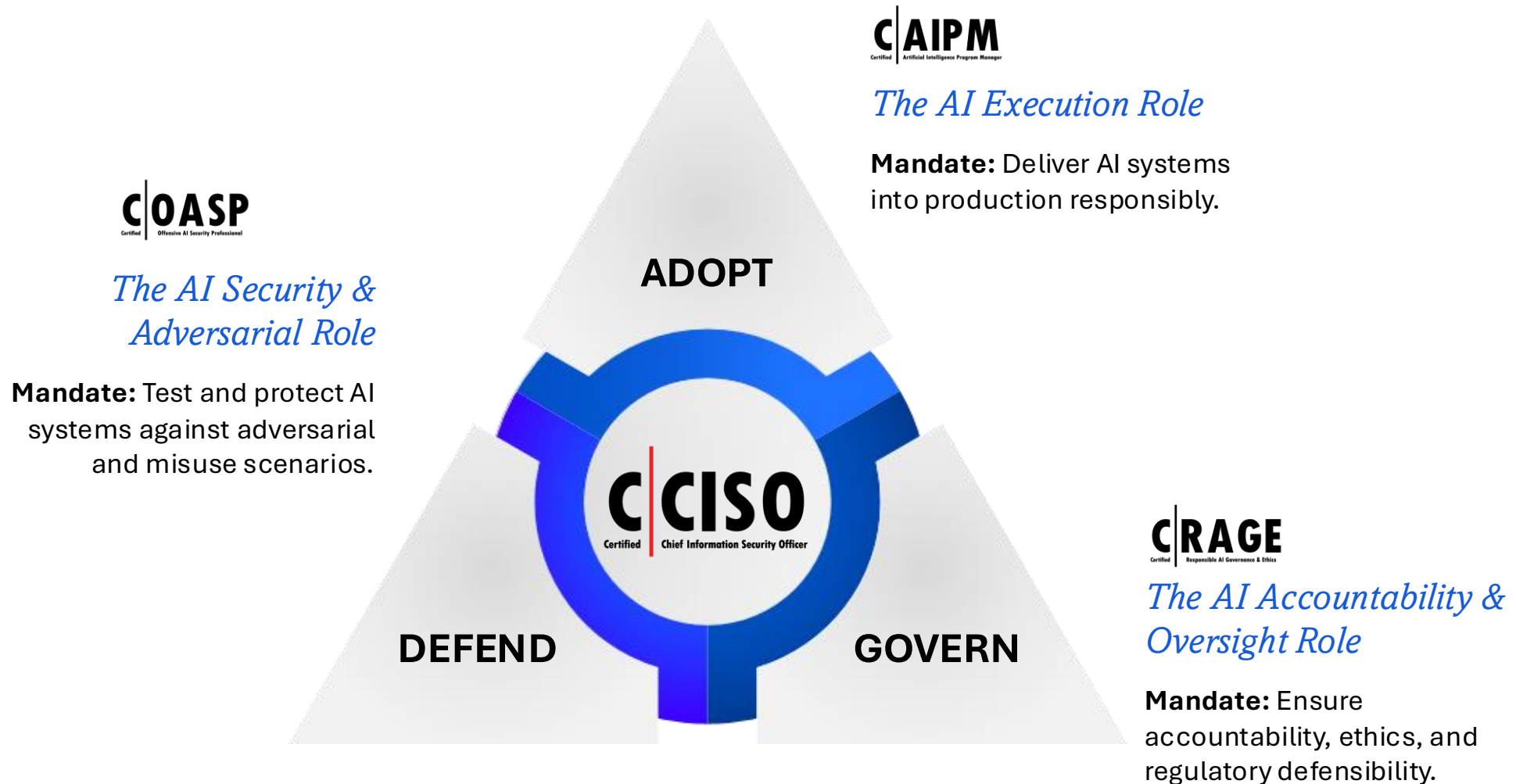


Regulated Industries

Aerospace, automotive (autonomous systems), and energy, where safety and compliance are non-negotiable.

Our Solution:

EC-Council's Proprietary ADG Framework For AI



D

E

F

E

N

D

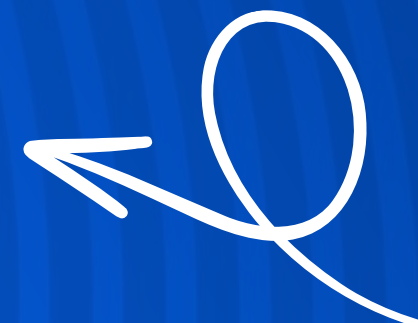
Introducing the
World's First



Certified Offensive AI Security Professional

Emulate. Exploit. Defend.

A **D** G Framework



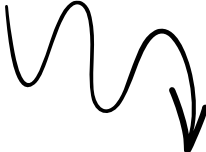


Who is a Certified Offensive *AI Security Professional?*

An Offensive AI Security Professional is a cybersecurity professional trained to think like an attacker and defend like an engineer. They probe AI systems from input to output, uncover hidden weaknesses, and design countermeasures so that AI deployments remain safe, reliable, and compliant.



C|OASP: *Professional Skills Framework*



Module 1	Offensive AI and AI System Hacking Methodology	}	Foundations & Readiness
Module 2	AI Reconnaissance and Attack Surface Mapping		
Module 3	AI Vulnerability Scanning and Fuzzing	}	Threat Strategy & Planning
Module 4	Prompt Injection and LLM Application Attacks		
Module 5	Adversarial Machine Learning and Model Privacy Attacks		
Module 6	Data and Training Pipeline Attacks	}	Attack Execution & Enablement
Module 7	Agentic AI and Model-to-Model Attacks		
Module 8	AI Infrastructure and Supply Chain Attacks		
Module 9	AI Security Testing, Evaluation, and Hardening	}	Defense Validation & Sustainment
Module 10	AI Incident Response and Forensics		



7 Key Features That Sets C|OASP Apart

01**150+ Hands-On Labs with Real AI Systems**

Attack and defend live LLMs including ChatGPT, Claude, and enterprise AI deployments in controlled lab environments.

02**DCWF-Aligned Learning Paths**

Curriculum mapped to DoD Cyber Workforce Framework for recognized career progression and government roles.

03**OWASP LLM Top 10 Coverage**

Master all 10 critical LLM vulnerabilities including prompt injection, insecure output handling, and training data poisoning.

04**Enterprise-Grade Attack Simulations**

Practice on realistic enterprise scenarios including RAG systems, AI agents, and multi-model architectures.

05**Red Team & Blue Team Perspectives**

Learn both offensive techniques to exploit AI systems and defensive strategies to protect them.

06**Certification Exam Included**

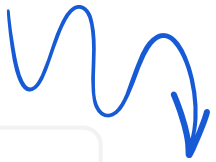
EC-Council official certification exam voucher included with your enrollment—no additional fees.

07**Access to AI Security Community**

Join a global network of AI security professionals, researchers, and practitioners for ongoing learning.



LLM Attack Tools & *Techniques Covered*



Reconnaissance

Enumeration Techniques
Endpoint Fingerprinting
Model Probing System
Prompt Extraction



Prompt Attacks

Prompt Injection
Context Overflow Exploits
Jailbreaking Techniques
Guardrail Evasion



Output Exploitation

XSS via LLM Outputs
SSRF Attacks Remote Code
Execution Trust Boundary
Bypass



Agent Attacks

Memory Manipulation
Task Flow Hijacking
Privilege Escalation Tool
Invocation Abuse



Supply Chain

Malicious Plugin
Development Typosquatting
Version Downgrade Attacks
Dependency Shadowing



Training-time Attacks

Data Poisoning Campaigns
Poisoned RAG Embeddings
Adversarial Perturbations
MLOps Lifecycle Threats



Model Attacks

Query-based Extraction
Model Fingerprinting
Surrogate Model Building
Weight Theft Techniques



Defensive Tools

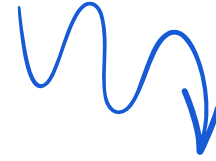
Input/Output Filtering
Sandboxing & Least Privilege
Drift Monitoring SIEM/SOAR
Integration



Frameworks & Standards

NIST AI RMF, ISO 42001,
EU AI Act, DoD DCWF
Mapping

Key *Features*



Training

Skill-Based Hands-On Labs:

Every module is anchored in practical exercises (prompt injection, model extraction, red vs. blue scenarios).

DCWF-Aligned Learning Path:

Each lab maps to specific DoD/NICE DCWF KSAs (e.g., 7027 reconnaissance, 7032 prompt injection, 7045 model extraction).

Full Adversarial Lifecycle:

Students progress from reconnaissance and prompt exploits to memory manipulation, supply-chain compromise, and AI risk reporting.

Top Skills

Neutralize AI-Native Cyberattacks Before They Spread.

Master reconnaissance, prompt injection, model theft, and full-chain adversarial testing to spot and shut down exploits unique to AI systems

Design AI Systems That Resist Adversaries by Default.

Apply defensive engineering, risk assessment, and compliance frameworks so that every AI deployment, from LLMs to multi-agent architectures. Is secure and auditable from day one.

Lead The Future Of AI Cyber Defense.

Step into next-generation roles like AI Red-Team Lead or Adversarial ML Architect, setting the standards for how enterprises and critical infrastructure defend against evolving AI threats.





Why Are Offensive AI Security Professional *Skills Essential for Any Cybersecurity Professional?*

AI is no longer just software; it's a dynamic, self-learning system. New exploits appear every week. Professionals who can test, break, and secure AI will be the ones every organization relies on. Traditional cybersecurity tools cannot detect adversarial AI exploits such as:



Prompt injections and context overflow



Model theft and surrogate model creation



Data poisoning during AI training



Multi-agent orchestration hijacking

Professionals need new methods and skills to evaluate, attack, and secure AI.

C|OASP-Enabled *Career Paths and Compensation*

Adversarial AI Security
Analyst

\$146K – \$230K

Adversarial Machine Learning
Researcher

\$125K – \$221K

AI Threat Hunter / AI Security
Analyst

\$113.6K – \$153.5K



AI Incident Response **Engineer** | \$103K – \$174K

Secure AI **Engineer / AI Security Architect** | \$179K – \$289K

AI Model Risk **Specialist** | \$105K – \$176K

AI Risk & Assurance **Specialist** | \$105K – \$176K

Security Program **Manager (AI Security)** | \$146K – \$220K

AI Product Security **Manager** | \$151K – \$239K

Real AI Job roles mapped to *Certified Offensive AI Security Professional* listed job vacancies In Top Companies

Engineering Analyst, Content Adversarial Red Team

 Google  Dublin, Ireland  Advanced

As a senior member of the team, you will mentor analysts, fostering a culture of continuous learning and sharing your expertise in adversarial techniques. You will also represent Google's AI safety efforts in external forums, collaborating with industry partners to develop best practices for responsible AI and solidifying our position as a thought leader in the field.

At Google we work hard to earn our users' trust every day. Trust & Safety is Google's team of abuse fighting and user trust experts working daily to make the internet a safer place. We partner with teams across Google to deliver bold solutions in abuse areas such as malware, spam and account hijacking. A team of Analysts, Policy Specialists, Engineers, and Program Managers, we work to reduce risk and fight abuse across all of Google's products, protecting our users, advertisers, and publishers across the globe in over 40 languages.

Responsibilities

- Design, develop, and oversee the execution of innovative red teaming strategies to uncover content abuse risks. Create and refine new red teaming methodologies, strategies and tactics.
- Influence across Product, Engineering, Research and Policy to drive the implementation of safety initiatives. Be a key advisor to executive leadership on content safety issues, providing actionable insights and recommendations.
- Mentor and guide junior and senior analysts, fostering excellence and continuous learning within the team. Act as a subject matter expert, sharing knowledge of adversarial and red teaming techniques, and risk mitigation.
- Represent Google's AI safety efforts in external forums and conferences. Contribute to the development of industry-wide best practices for responsible AI development.
- Be comfortable to be exposed to graphic, controversial or upsetting content.

C | OASP course matches ~93% of the *skills* required for below

Google – Engineering Analyst, Content Adversarial Red Team

JD of Google – *Engineering Analyst, Content Adversarial Red Team*

Description: “In this pivotal role, you will come with considerable direct experience in adversarial testing and red teaming, particularly of Generative AI, so that you design and direct red teaming operations, creating innovative methodologies to uncover novel content abuse risks.”

Google

% Overlap with *COASP* Skills

93% Full-scope adversarial red-teaming of generative AI, abuse risk discovery, innovative testing.

Deloitte.

If you are a technology visionary with a passion for transforming global tax business with digital technology, consider working with the US Tax Transformation technology team. This is an exciting opportunity to support global execution of Deloitte's tax strategy as we shift from "doing digital" to "being digital" by reimagining how we engage with our clients, deliver our services, operate our business, and create value.

Work you'll do

As a Deloitte Lead AI Security Engineer, you will be crucial in safeguarding our advanced AI models, data, and infrastructure. You'll work closely with Data Scientists, Data Engineers, and MLOps/DevOps teams.

Additional responsibilities include:

- Implement defences against AI-specific attacks (adversarial, prompt injection, data leakage)
- Conduct AI-focused security assessments, penetration tests, red/purple team exercises
- Analyse AI system vulnerabilities, develop mitigation strategies, and create AI risk heat maps
- Implement security controls throughout the AI/ML lifecycle (data handling, training with GPU isolation, deployment, monitoring, versioning, provenance). Integrate SAST/DAST for ML artifacts
- Manage audit trails and automated compliance checks
- Implement AI-specific incident response and develop regulatory disclosure playbooks
- Manage AI security monitoring, implement executive dashboards linking security to business KPIs, develop security metrics (Adversarial Risk Score, Model Drift Index)
- Implement secure training environments and fine-grained data access controls
- Contribute to AI-generated fraud detection in transaction monitoring systems.
- Act as an AI security SME, continuously research emerging threats

C | OASP course matches ~88% of the skills required for below

Deloitte – Lead AI Security Engineer

JD of Deloitte – *Lead AI Security Engineer*

The job description mentions “expertise in identifying and mitigating AI/ML security threats, including adversarial attacks, prompt injection, and data leakage.”

Also states familiarity with security frameworks (e.g., NIST AI RMF, ISO 27001) and ability to apply them in enterprise AI environments

% Overlap with COASP Skills

88% Strong alignment with adversarial testing, LLM hardening, NIST/ISO/EU AI Act compliance, and SOC/SIEM integration.



CAREERS HOME LOCATIONS ▼ EARLY CAREER

Things You'll Do

- Design and architect end-to-end security for the entire ML/AI pipeline, developing reference architectures for various deployment patterns, and ensuring compliance
- Perform threat modeling for AI/ML systems, identifying vulnerabilities in data poisoning, model evasion, extraction, and integrity; define and implement data privacy, DLP, and encryption standards for sensitive training and inference data; and collaborate with Data Scientists and ML Engineers to integrate security-by-design principles from initial concept.
- Implement and monitor technical controls for model integrity, governance, and AI security, including access control and bias detection.
- Ensure AI systems comply with industry standards, privacy regulations (e.g., GDPR, CCPA), and internal Responsible AI policies, working with Legal and Compliance to document and validate technical adherence to AI governance frameworks.

C | OASP course matches ~82% of the skills required for below

Qualtrics – AI Security Architect

JD of Qualtrics – *AI Security Architect*

The job: “Essential to designing, building, and maintaining robust security architectures for our AI.

While the posted snippet is short, the title and context suggests architecture of AI security, which implies many of your skills around system hardening, secure context windows, plugin security, agent architectures.

% Overlap with *COASP* Skills

82% Solid on secure AI architecture, sandboxing, input/output filtering, and resilience design.

Adversarial AI Engineer

GoFundMe

San Francisco, CA

• Adversarial Testing & Red Teaming

- Execute adversarial testing of LLMs, Agentic AI systems, recommendation models, and fraud detection tools using techniques such as prompt injection, jailbreaking, data poisoning, model inversion, and membership inference attacks.
- Develop synthetic attack datasets tailored to fundraising and trust & safety scenarios.

• Build Security Frameworks & Defenses

- Develop automated adversarial testing pipelines integrated into CI/CD, create reusable robustness evaluation libraries, and generate synthetic attack datasets for fundraising scenarios.
- Build reusable robustness evaluation libraries.
- Deploy real-time detection for prompt injection and model evasion, implement input validation, output filtering, adversarial training, and differential privacy mechanisms.

• Cross-Functional Security Leadership

- Partner with Trust & Safety on fraud countermeasures.
- Collaborate with Product Security on AI threat models and governance.
- Work with Data Science to mitigate algorithmic bias and ensure robust defense.

• Governance & Policy

- Establish AI security policies, training, and deployment review processes aligned with NIST AI RMF.
- Build monitoring and incident response systems for AI security.
- Stay current with emerging attack vectors and defense mechanisms.
- Contribute to open-source adversarial tools.
- Publish and speak externally to advance GoFundMe's leadership in AI security.

C | OASP course matches ~82% of the skills required for the

GoFundMe – Adversarial AI Engineer Role

JD of GoFundMe - *Adversarial AI Engineer*

We're looking for an Adversarial AI Engineer who combines offensive security expertise with machine learning depth. You will execute red-team testing against our AI systems, build automated adversarial evaluation pipelines, and deploy defense mechanisms to keep our AI safe.

% Overlap with *COASP* Skills

95% Near-perfect fit: adversarial red-team testing, exploit simulation, automated evaluation pipelines, AI defense mechanisms.

Job Roles Mapped to C|OASP

**AI Red Team
Specialist/Adversarial AI Engineer**

**Offensive Security
Engineer (AI/LLM Focus)**

**Adversarial AI
Security Analyst**

**Adversarial Machine
Learning Researcher**

**AI Threat Hunter/AI
Security Analyst**

**AI Malware &
Exploit Analyst**

**AI Incident
Response Engineer**

**AI Test & Evaluation
Specialist**

**Secure AI Engineer/AI
Security Architect**

**MLOps/AI Ops
Security Specialist**

**AI Model Risk
Specialist**

**LLM Systems
Engineer**

**AI Risk & Assurance
Specialist**

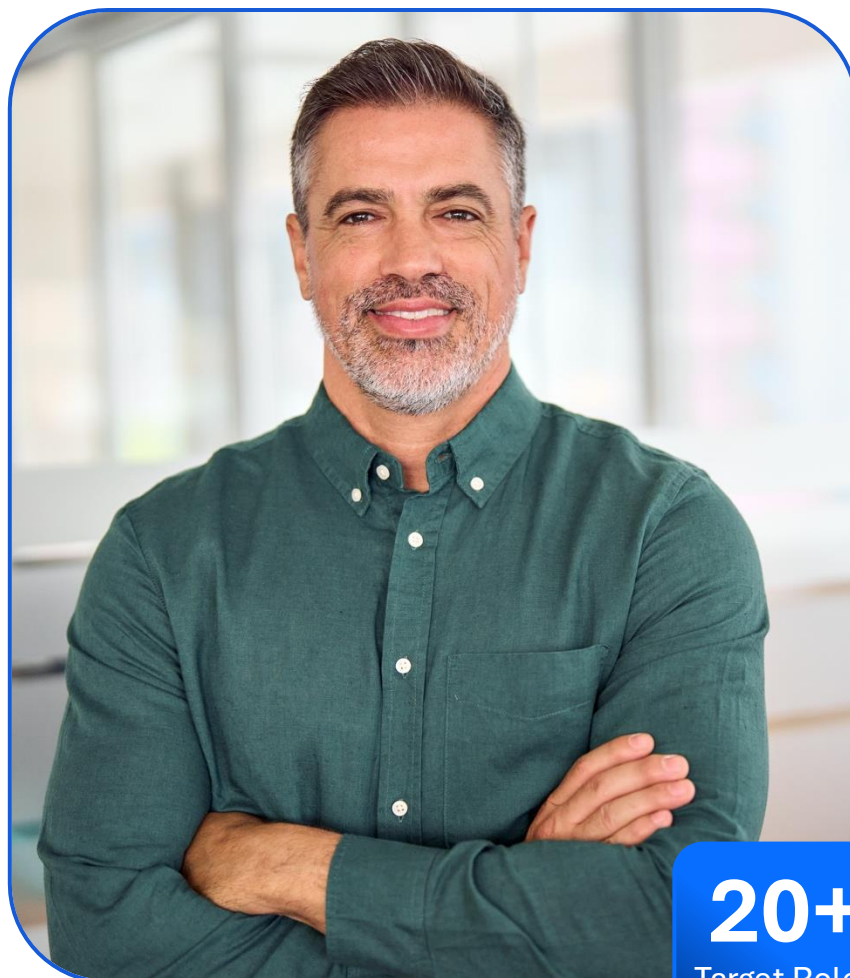
**Cyber Threat Intelligence
(CTI) Analyst – AI Focus**

**Security Program
Manager (AI Security)**

**AI Product Security
Manager**

**AI Risk
Advisor/Consultant**

C|OASP : *This Program Is Ideal For*



20+
Target Roles

C|OASP is designed for security professionals who want to master offensive and defensive AI security techniques.

OFFENSIVE SECURITY

- Penetration Tester/Ethical Hacker
- Red Team Operator/Red Team Lead
- Offensive Security Engineer
- Adversary Emulation/Purple Team Specialist

DEFENSIVE SECURITY

- SOC Analyst (Tier 2/3)/Detection Engineer
- Blue Team Engineer/Threat Detection Engineer
- Incident Responder (IR)/DFIR Analyst
- Security Operations Manager (SOC Lead)

THREAT INTELLIGENCE

- Malware Analyst/Threat Researcher
- Cyber Threat Intelligence (CTI) Analyst – AI Focus
- Fraud/Abuse Detection Analyst (AI-enabled threats)

AI/ML ENGINEERING

- ML Engineer/Applied AI Engineer
- GenAI Engineer (RAG/Agents)
- AI/LLM Application Developer
- MLOps/AI Platform Engineer

SECURITY ENGINEERING

- DevSecOps/Secure DevOps Specialist
- Application Security Engineer (LLM Apps/APIs)
- Product Security Engineer/AI Product Security

AI SECURITY ARCHITECTURE

- Secure AI Engineer/AI Security Architect
- LLM Systems Engineer

Persona Profile

For C|OASP



The Red Team Veteran – *Steve, the Ethical Hacker* *Evolving with AI*

Age: 28–38

Current Role: Penetration Tester/Red Team Lead

Experience: 5–10 years in offensive security

Goal: Stay ahead of AI-driven defenses and attackers.

Motivation:

- Become one of the first in the world to **master offensive AI exploitation**.
- Build a portfolio of adversarial AI projects to pitch elite employers.

Key Messaging Angle:

“You’ve mastered human hacking. Now learn how to hack AI itself.”

Upgrade from Red Team to Adversarial AI Operator.

Pain Points:

- Traditional pen testing feels outdated against AI-powered defenses.
- Struggles to simulate adversaries that use LLMs or AI agents.
- Wants to differentiate from “just another C|EH” or CPENT-level tester.

How COASP Helps:

- Teaches offensive use of AI to simulate real adversarial AI attacks.
- Expands red team skill set into model extraction, plugin abuse, and prompt injection.
- Builds credibility with employers seeking *AI-augmented pen testers*.

Targeting:

Online: LinkedIn Groups (Pentesters, Red Team Ops), Reddit (r/netsec, r/redteamsec), YouTube hacking influencers, TryHackMe/HTB communities.

Offline: Cyber ranges, hacking conferences (DEF CON, Black Hat, c0c0n, Nullcon), university hacking clubs.



The SOC Sentinel – *Sara, the Blue Teamer Who Predicts the Future*

Age: 26–36

Current Role: SOC Analyst/Blue Team Engineer

Experience: 3–8 years in monitoring and threat detection

Goal: Detect and defend against AI-assisted cyberattacks.

Motivation:

Integrate **AI threat detection** into existing SOC pipelines.

Move toward **AI Security Analyst** or **Adversarial Defense Engineer** roles.

Key Messaging Angle:

“Defend not just your systems, but your AI itself.”

Transform your SOC from reactive to AI-aware.

Pain Points:

Traditional SIEM/SOAR playbooks don’t detect AI-generated threats.

Feels underprepared for AI-integrated incidents or model poisoning alerts.

How C|OASP Helps:

Enables SOC teams to detect and respond to AI-specific anomalies.

Adds AI event response automation in SIEM/SOAR systems.

Bridges the gap between traditional monitoring and AI security.

Targeting:

Online: LinkedIn cyber defense and SOC analyst groups, SIEM vendor communities (Splunk, IBM QRadar, Microsoft Sentinel).

Offline: Corporate blue team workshops, SOC manager training programs, and regional CERT or ISAC meetups.





The Threat Hunter – *Diego, the Researcher Who Hunts the Unknown*

Age: 25–35

Current Role: Threat Hunter/Cyber Threat Intelligence Analyst

Experience: 3–6 years

Goal: Hunt AI-based TTPs and predict adversarial models before they hit production.

Motivation:

Stay two steps ahead of attackers experimenting with adversarial AI.

Publish or present research on **AI-driven attack surfaces**.

Key Messaging Angle:

“AI adversaries are already hunting you—it’s time you hunt them first.”

Lead the next wave of cyber threat intelligence.

Pain Points:

New AI attack chains (poisoning, model extraction, data leakage) are invisible to standard tools.

Hard to analyze adversarial datasets or synthetic payloads.

How C|OASP Helps:

Helps track AI-assisted threat actors and emerging adversarial TTPs.

Teaches data poisoning and evasion pattern recognition.

Enables proactive defense with predictive adversarial models.

Targeting:

Online: Threat intel blogs (Mandiant, MITRE ATT&CK), X (Twitter) infosec handles, and Discord/Slack threat-hunting communities.

Offline: Cyber threat hunting bootcamps, ISACA/ISC² chapter events, and cyber range labs.



The AppSec Guardian – *Meera, the Developer Who Builds Secure AI Systems*

Age: 25–40

Current Role: Application Security Engineer/Secure DevOps Specialist

Experience: 4–10 years

Goal: Protect LLM-integrated apps and APIs from prompt injection and data leaks.

Motivation:

Learn **sandboxing, input/output filtering**, and **AI-specific secure coding**.

Become a go-to expert for **AI-secure application development**.

Key Messaging Angle:

“You don’t just code apps—you build trust in AI.”

Secure LLMs before attackers exploit them.

Pain Points:

Struggles to secure generative AI features within web apps or chatbots.

No structured framework for testing or hardening AI components.

How C|OASP Helps:

Teaches secure design for AI-integrated apps and APIs.

Covers sandboxing, filtering, and context hardening for LLM systems.

Helps prevent prompt injection and plugin manipulation in production apps.

Targeting:

Online: GitHub security repositories, OWASP forums, Stack Overflow DevSecOps threads, and LinkedIn AppSec pages.

Offline: DevSecOps meetups, secure coding workshops, and AppSec conferences.





The Forensic Responder – *Lina, the IR Specialist* *Who Investigates AI Breaches*

Age: 27–38

Current Role: Incident Responder/Digital Forensics Analyst

Experience: 4–8 years

Goal: Understand how AI-generated attacks manifest and how to respond.

Motivation:

Build capabilities for **AI breach response** and **adversarial forensics**.
Prepare for roles in **AI CERT teams** or **AI-specific IR units**.

Key Messaging Angle:

“Tomorrow’s incidents won’t look like yesterday’s breaches.”
Master AI incident response before the first AI breach hits your network.

Pain Points:

Unclear evidence chains in AI-led or AI-assisted breaches.
No tools or methods for AI-specific forensic investigation.

How C|OASP Helps:

Trains IR teams to analyze AI-led or AI-assisted breaches.
Introduces frameworks for adversarial forensics and AI response playbooks.
Enables rapid containment of AI-induced misbehavior in systems.

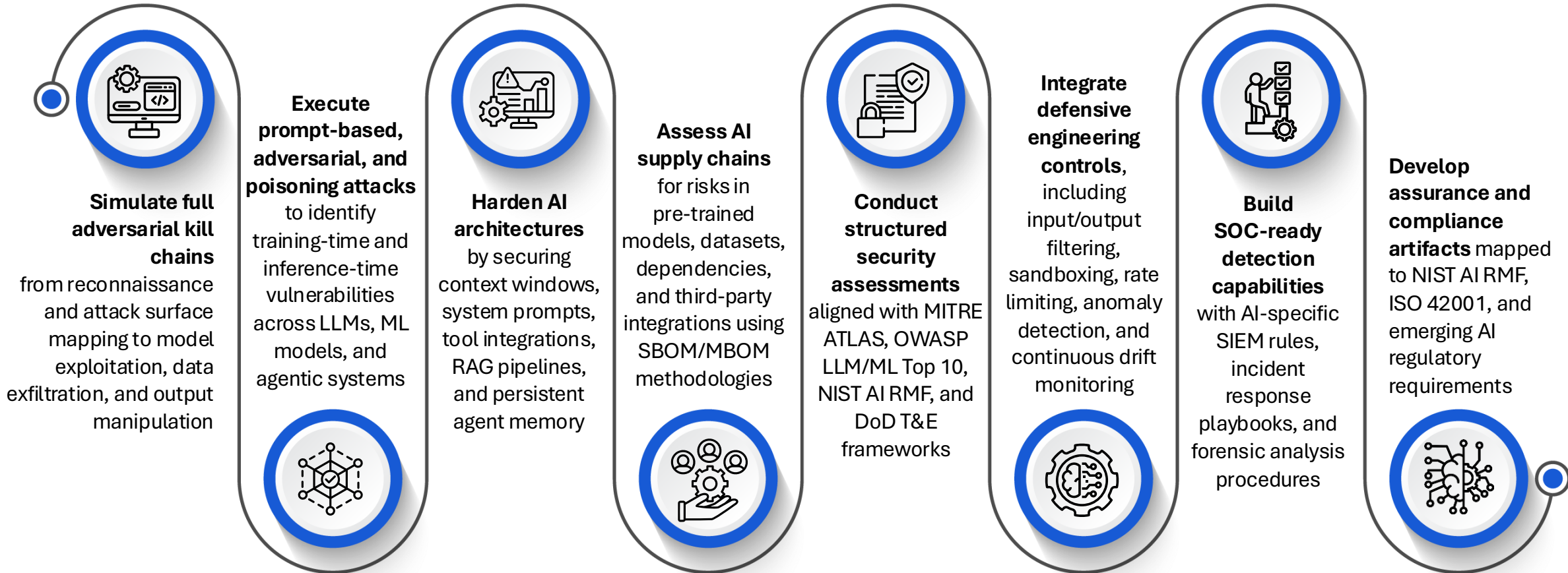
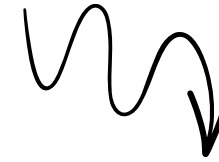
Targeting:

Online: DFIR Discord channels, SANS DFIR blogs, and LinkedIn IR communities.

Offline: Digital forensics summits, IR tabletop exercises, and national CERT events.



Key Learning *Outcomes from the Program*

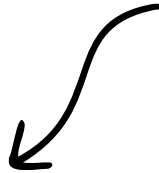


C|OASP vs. *Market*

Capability / Outcome Area	COASP (EC-Council)	SANS - SEC535	OffSec - LLM Red Teaming	Hack The Box - AI Red Teamer	TONEX - COAIS
Offensive AI and AI System Hacking Methodology	✓	●	●	✓	●
AI Reconnaissance and Attack Surface Mapping	✓	✓	●	●	●
AI Vulnerability Scanning and Fuzzing	✓	●	✗	✗	✗
Prompt Injection and LLM Application Attacks	✓	✗	✓	✓	✓
Adversarial Machine Learning & Model Privacy Attacks	✓	✗	✗	✓	●
Data and Training Pipeline Attacks	✓	✗	●	✓	●
Agentic AI & Model-to-Model Attacks	✓	●	●	●	●
AI Infrastructure & Supply Chain Attacks	✓	✗	✓	●	●
AI Security Testing, Evaluation, and Hardening	✓	✗	●	●	✗
AI Incident Response and Forensics	✓	✗	✗	✗	✗

Note: This is EC-Council's view of our competitive landscape.

Training and *Exam Information*



Training *Details*

Title of the Course: Certified Offensive AI Security Professional (C | OASP)

Version: 1

Training Duration: 5 days (40 Hours)

Training Prerequisite: 3 Years of Cybersecurity Experience

Delivery Mode:

- Instructor-led Training (ILT)
- iWeek (Synchronous Online Learning)
- iLearn (Asynchronous Online Learning)

Exam *Details*

Exam Title: Certified Offensive AI Security Professional

Exam Code: 312-52

Number of Questions: 70 (65 MCQ + 5 PBQ)

Duration: 6 hours

Availability: ECC Exam Portal

Passing Score: 70-80%

Test Format: Multiple Choice Question & Performance Based Questions



Thank You!

