



Securing AI—Delivered through Cognizant Secure AI Services and CrowdStrike

Overview

AI is rapidly becoming central to enterprise strategy, powering automation, decision intelligence and personalized experiences at scale. Adoption is accelerating at an unprecedented pace, global AI spending is projected to reach \$632 billion by 2028 according to IDC,¹ while enterprises are already putting 11 times more AI models into production year over year as reported by Databricks.² Gartner further predicts that 80% of generative AI business applications will be built on enterprise data platforms by 2028, underscoring how deeply AI will be embedded into core systems and workflows.³

As organizations scale AI, they increasingly depend on third-party LLMs, API-driven AI services and Model Context Protocol (MCP) servers. While these accelerate innovation, they also expand the attack surface introducing risks such as data leakage, model inversion, unauthorized retention, supply chain compromise and unpredictable model behavior. Enterprises must also defend against AI-specific threats including prompt injection, data poisoning, adversarial manipulation and jailbreak attacks.

Key challenges

Securing AI now requires more than traditional infrastructure controls. Traditional security was built for deterministic software and release-gate controls. AI systems are probabilistic and context driven. They learn, drift and depend on enterprise context that can be poisoned, manipulated or exfiltrated. As agents move into core workflows, a single compromised context or prompt can trigger confidently wrong actions at scale—creating systemic risk that legacy, fragmented security wasn't designed to prevent. As organizations scale AI adoption, CISOs must address several emerging challenges:

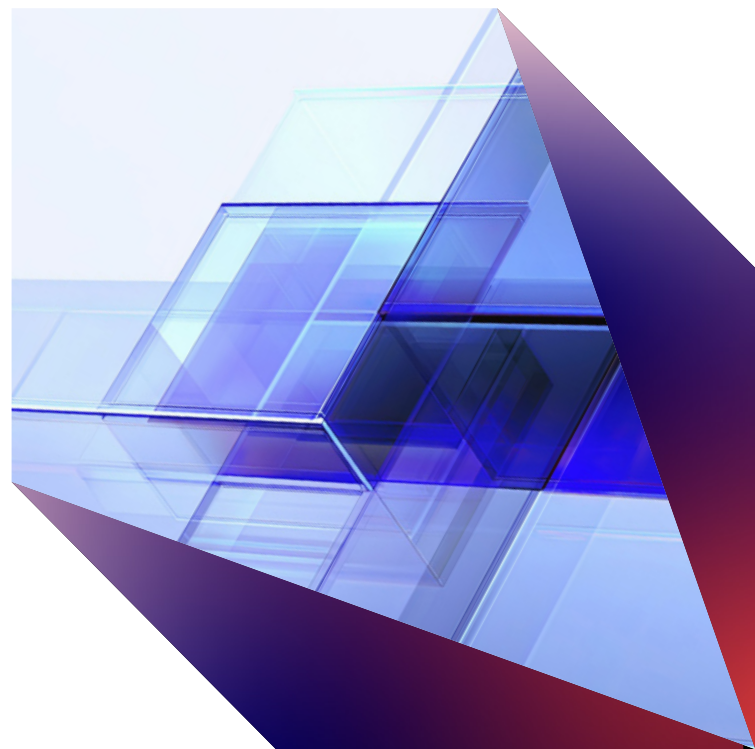
- Rapid growth of enterprise AI applications is significantly expanding the organizational attack surface
- New AI-specific threats such as prompt injection, data poisoning, model tampering and agent impersonation are emerging
- Autonomous and agentic AI systems increase risk through uncontrolled interactions with browsers, APIs, databases and tools
- AI is being actively weaponized for high-impact fraud, deepfake impersonation, phishing and credential theft
- Insecure APIs, credentials and misconfigured agents enable hijacking of AI services and unauthorized actions
- Lack of run-time visibility and monitoring makes stealthy AI-driven attacks difficult to detect and respond to
- AI security incidents pose significant business risk, including financial loss, compliance penalties and reputational damage
- The convergence of AI security and responsible AI is expanding the CISO mandate, requiring new skills, governance frameworks and organizational capabilities

If a model is secure but biased, it becomes a liability. If it is ethical but vulnerable to injection, it becomes a risk.

Securing the AI-powered enterprise

Organizations must adopt AI-specific governance, zero trust for model interactions, secure integration patterns for external LLM ecosystems and continuous monitoring of model behavior. Aligning with frameworks like NIST AI RMF and ISO/IEC 42001 helps ensure responsible and compliant adoption.

By embedding security across the entire AI lifecycle from data ingestion to deployment, and across the entire AI stack from AI infrastructure to AI applications and services, enterprises can safeguard intellectual property, deploying responsible and AI governance practices, maintain regulatory confidence and unlock the full strategic value of AI while mitigating emerging risks.



AI introduces a dual mandate for cybersecurity: accelerating defense with AI (AI for Security) while securing AI systems themselves (Cognizant® Secure AI Services). These mandates are inseparable; enterprises cannot protect the AI stack without the speed, consolidation and orchestration that AI-driven defense delivers. Strengthening both sides of this mandate is essential for realizing AI's full strategic value while safeguarding enterprise systems, data and operations.

Key foundation principles

We actively align with leading global AI standards and embed their principles into our lifecycle controls and governance practices. Our adaptive governance model provides continuous oversight and real-time assurance, so systems remain accountable, consistent and aligned as technologies and expectations evolve.

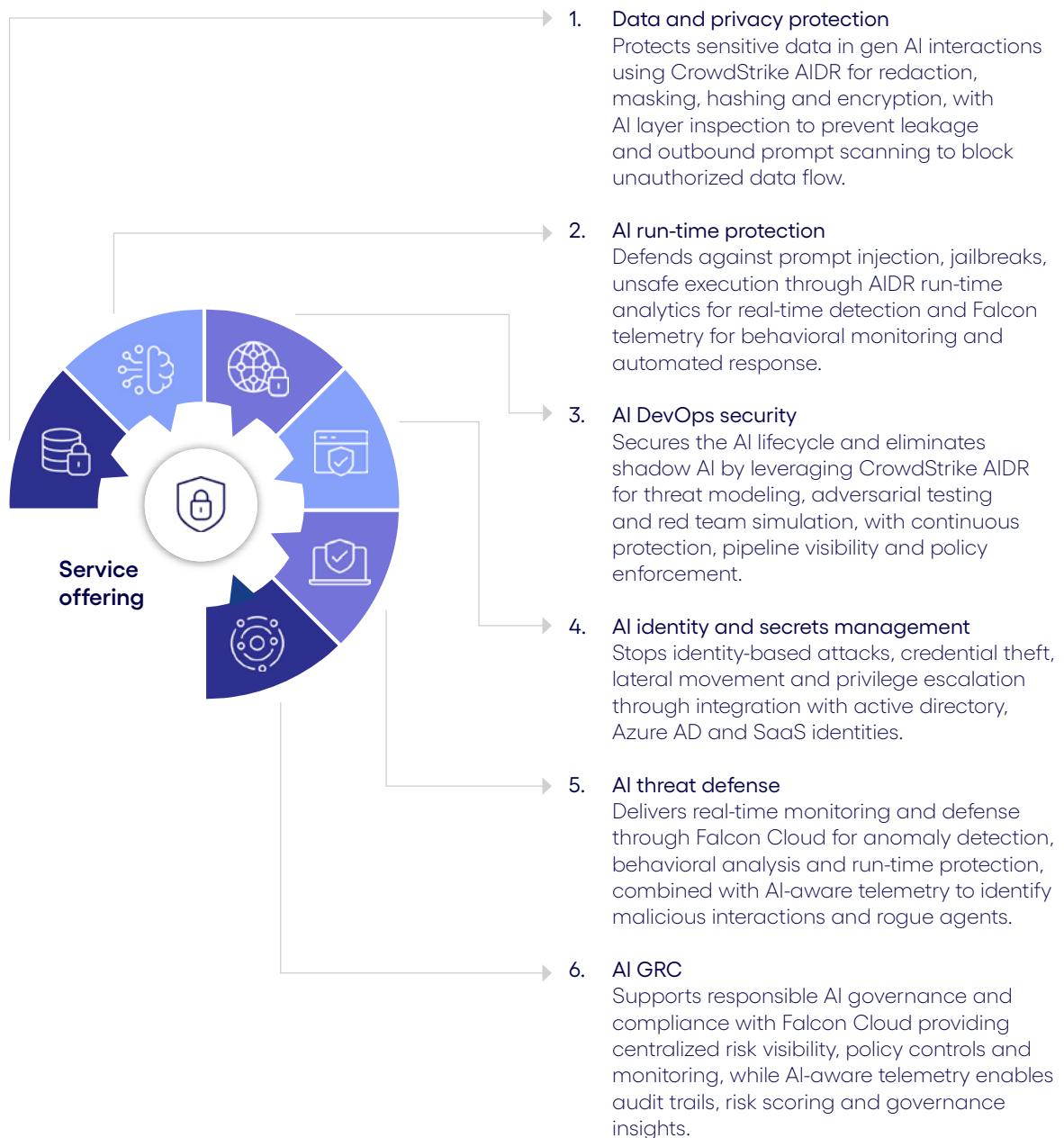
Our Secure AI Services governance model is built on three foundations: secure Agent Development Lifecycle (ADLC) (build-time assurance), Neuro® Cybersecurity (consolidated control plane) and responsible AI (trust and accountability).

Secure Agent Development Lifecycle	Neuro Cybersecurity	Responsible AI
Given that agentic systems are nondeterministic—they are autonomous, adaptive, goal-oriented and capable of proactive behavior—we address these challenges by providing a repeatable, outcome-driven framework that enables agents to be purpose-built, compliant and optimized for enterprise success. Our secure ADLC process is designed for creating, testing and monitoring nondeterministic, autonomous AI agents that reason and act. Providing security and governance integrated across design, build, test, deploy and change of agentic systems.	As threats become increasingly sophisticated and pervasive, organizations need to stay ahead of bad actors. Siloed operations and risk mitigation decisions often lag the need to act in real time, leading to higher risk of exposure. Need of the hour is AI-enabled, easy-to-consume UI that enables decisions across the enterprise. Cognizant Neuro Cybersecurity revolutionizes risk management, threat and vulnerability management and compliance assurance—by integrating and orchestrating point solutions to ensure comprehensive security coverage using AI. It is the “unification layer” for AI security posture + run-time signals + response actions: - Consolidates AI signals (model, prompt, RAG, agent actions) with enterprise signals (identity, cloud, network, endpoint) - Enables faster correlation and response orchestration across tools - Produces audit-ready evidence and posture dashboards for AI risk	Our focus is on operationalizing trust in AI systems through continuous assurance, control and accountability across the lifecycle. In addition to offer a compliance layer, responsible AI functions as a continuous trust layer that: - Ensures AI systems behave as intended in real-world conditions - Provides traceability of decisions, prompts and actions - Enables policy enforcement and behavioral guardrails for AI agents - Continuously validates performance, bias, drift and safety - Generates audit-ready evidence aligned to ISO 42001, NIST AI RMF and regulatory requirements Through Cognizant Trust™, our responsible AI assurance framework, organizations gain: - AI system inventory and risk classification - Continuous compliance monitoring and policy validation - Incident tracking and remediation workflows - Real-time dashboards and metrics for AI system performance and trust This enables organizations to move from static compliance to continuous, provable trust in AI systems.

Service offering

Cognizant leverages CrowdStrike Falcon® AI Detection and Response (AIDR) which redefines AI security with comprehensive protection for both employee adoption of generative AI (gen AI) tools and run-time security for homegrown AI development. Our solution powered by Falcon AIDR provides unified visibility, real-time threat detection, data protection, access controls and automated response across endpoints, applications, AI agents, MCP servers, AI/API gateways and cloud environments through a single sensor and unified console.

Cognizant's Secure AI Services approach provides comprehensive, end-to-end protection by strengthening the core pillars across the AI lifecycle: infrastructure security, data security, run-time security, AI DevOps security, identity protection, threat defense, and governance, risk and compliance (GRC) security integrated with CrowdStrike's AIDR solution, which is a single agent-based solution for unified visibility and response.



Secure AI Services delivers consistent governance, ethical guardrails and secure-by-design engineering practices across the design, development, deployment and operational stages of AI systems. Together, these capabilities enable organizations to protect models, data and interactions, reduce exposure to AI-specific threats and maintain trust and compliance at scale.

Engagement methodology

Establish clarity for informed decisions



Assess

- Posture discovery
- Benchmarking
- Risk profiling

- AI security posture assessment
- Risk identification and logging
- Regulatory and standards mapping
- Threat modeling for AI/ML
- Red teaming prep

Build scalable, resilient security



Architect

- Design
- Implement
- Integrate

- Secure AI architecture design
- Technology and platform selection
- Governance and RACI structure
- Red teaming integration
- Ethical AI compliance framework

Enable rapid threat and incident response



Act

- Monitor
- Detect
- Respond

- Deployment and onboarding
- Control testing and validation
- Incident response and breach readiness
- Red teaming execution
- Resilience and recovery planning

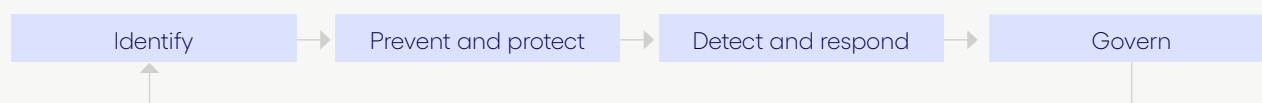
Ensure compliance and resilience



Assure

- Measure
- Validate
- Improve

- Continuous compliance and audit
- Security awareness and training
- Performance and capacity management
- Periodic red teaming testing
- KPI/SLA tracking and reporting



Use cases

- Prevent deepfake-driven fraud and impersonation
- Discover shadow AI across browsers, applications, AI/API gateways MCP servers and cloud environments
- Map AI relationships to understand AI usage patterns and identify risky behaviors
- Identify malicious entities and indicators of compromise (IOCs) and validate MCP interactions
- Detect harmful content and enforce content alignment for AI uses
- Protect sensitive data across AI interactions by preventing data leaks to third-party LLMs and redact sensitive information
- Defend against prompt injection and jailbreak attacks
- Secure AI-driven data pipelines
- Safeguard AI agents and autonomous workflows
- Secure AI supply chain and third-party dependencies
- Enable AI-specific threat detection and incident response
- Ensure regulatory compliance and audit readiness for AI

Why Cognizant?

Cognizant helps enterprises turn AI ambition into real, scalable business outcomes. Our approach goes beyond experimentation to embed AI deeply into core operations, decision-making and customer experiences. By combining deep industry expertise, proprietary AI platforms and responsible engineering, we deliver AI solutions that are contextual, secure and built for enterprise scale.

At the heart of this approach is Cognizant Neuro, our enterprise AI platform suite that acts as a “unification layer” for AI security posture, run-time signals and automated response actions, and accelerates the journey from identifying high-value use cases to deploying and governing AI and agentic solutions in production. Our partnership with CrowdStrike enhances this capability by integrating industry-leading AI-aware security technologies from the Falcon platform, including Falcon AIDR, Falcon Next-Gen SIEM, Falcon Identity Protection and Falcon Cloud Security. Together, these capabilities deliver seamless policy enforcement, continuous monitoring and real-time threat detection across AI applications, ensuring AI is adopted responsibly and securely.

Cognizant’s human-centred AI philosophy ensures technology augments human intelligence, drives productivity and enables new ways of working. Supported by CrowdStrike and a platform-agnostic partner ecosystem, we deliver future-ready AI solutions without vendor lock-in. With a relentless focus on outcomes, trust and scale, Cognizant enables organizations to unlock sustained value from AI responsibly, securely and at speed.

To know more, please contact: cybersecuritypulse@cognizant.com

Sources

1. Worldwide Spending on Artificial Intelligence Forecast to Reach \$632 Billion in 2028, According to a New IDC Spending Guide
2. State of AI: Enterprise Adoption & Growth Trends | Databricks Blog
3. Gartner Predicts by 2028, 80% of GenAI Business Apps Will Be Developed on Existing Data Management Platforms



Cognizant (Nasdaq: CTSH) is an AI Builder and technology services provider, bridging the gap between AI investment and enterprise value by building full-stack AI solutions for our clients. Our deep industry, process and engineering expertise enables us to build an organization's unique context into technology systems that amplify human potential, drive tangible outcomes and keep global enterprises ahead in a fast-changing world. See how at cognizant.ai or @cognizant.

World Headquarters

300 Frank W. Burr Blvd.
Suite 36, 6th Floor
Teaneck, NJ 07666 USA
Phone: +1 201 801 0233
Fax: +1 201 801 0243
Toll Free: +1 888 937 3277

European Headquarters

280 Bishopsgate
London
EC2M 4RB
England
Tel: +44 (0)1 020 7297 7600

India Operations Headquarters

5/535, Okkiam Thorapakkam,
Old Mahabalipuram Road,
Chennai 600 096
Tel: 1-800-208-6999
Fax: +91 (0) 44 4209 6060

APAC Headquarters

1 Fusionopolis Link, Level 5
NEXUS@One-North, North Tower
Singapore 138542
Tel: +65 6812 4000

© Copyright 2025–2027, Cognizant. All rights reserved. No part of this document may be reproduced, stored in a retrieval system, transmitted in any form or by any means, electronic, mechanical, photocopying, recording or otherwise, without the express written permission of Cognizant. The information contained herein is subject to change without notice. All other trademarks mentioned herein are the property of their respective owners.