

2025

AI Threats & Market Insights Enterprise CISOs

1. Executive Insights: The Adversary's New Playbook

1

The modern adversary doesn't look like yesterday's lone hacker. They now operate like well-funded enterprises: outsourcing initial access, productizing intrusion services, and scaling globally with the help of AI.

- **Speed is the new weapon:** Average breakout time — the window between initial compromise and lateral movement — fell to just 48 minutes, with the fastest case at 51 seconds. That means defenders may have less than a coffee break to detect and respond.
- **Malware is optional:** 79% of observed intrusions required no malware, relying instead on valid credentials, remote management tools, and social engineering. Traditional signature-based defenses are increasingly irrelevant.
- **Voice phishing becomes mainstream:** Vishing campaigns grew by 442% in 2024, using AI voice cloning to impersonate IT staff and executives with unnerving accuracy.
- **Nation-states accelerate operations:** China-linked activity rose 150% across all sectors, with financial services, manufacturing, and media facing 200–300% more intrusions than the prior year. DPRK's FAMOUS CHOLLIMA ran malicious insider schemes at scale, with 40% of its operations insider-driven.
- **The human layer is cracking:** Social engineering, insider placement, and identity abuse are now the primary gateways — not exploits.

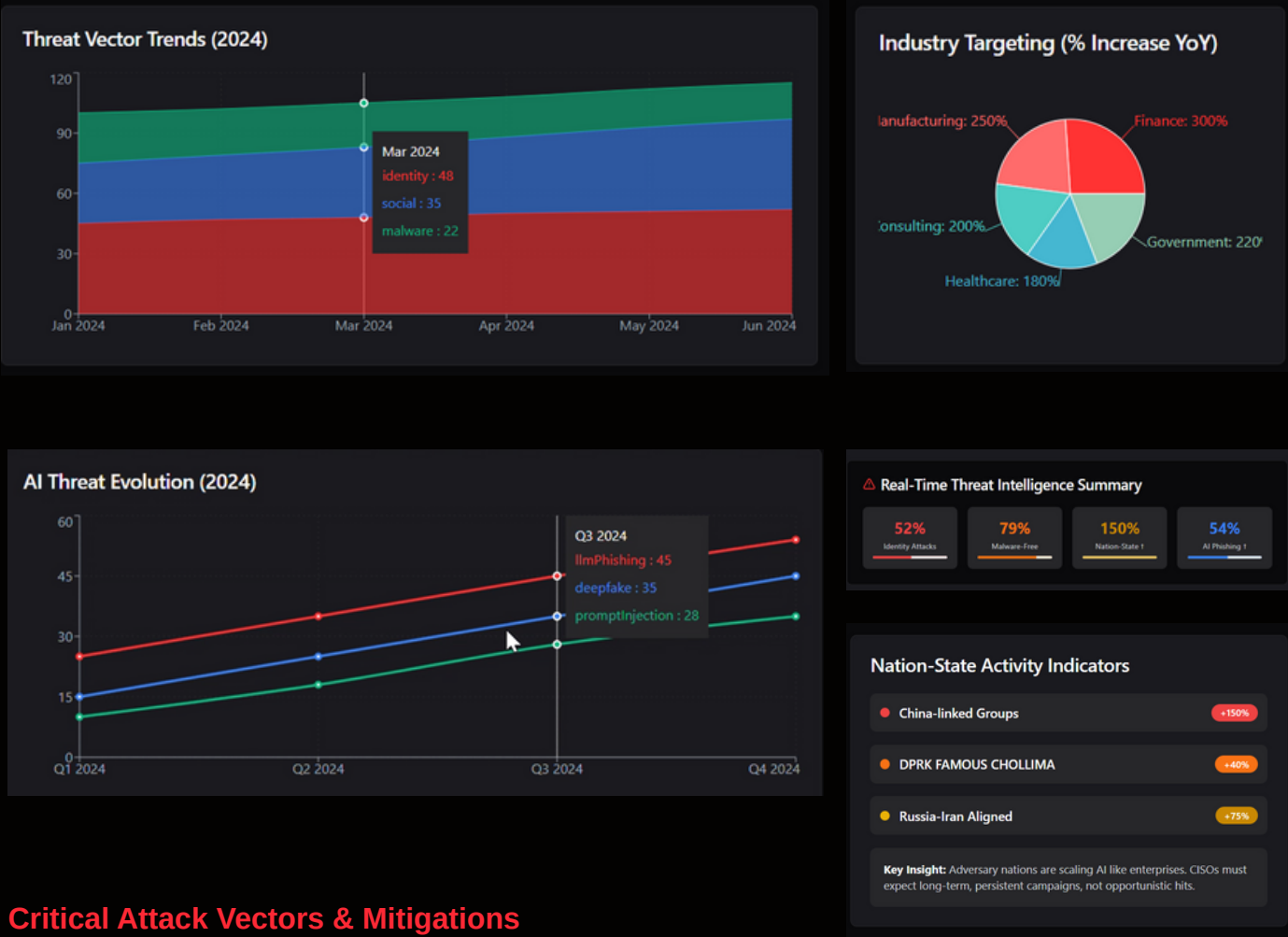
CISO Implication: Defenses built for yesterday's malware-centric world are too slow and too narrow.

Adversaries are innovating like enterprises, forcing security leaders to adopt enterprise-grade controls that close identity gaps, monitor insider risks, and cut response windows to near-real time.

CISO Business Intelligence

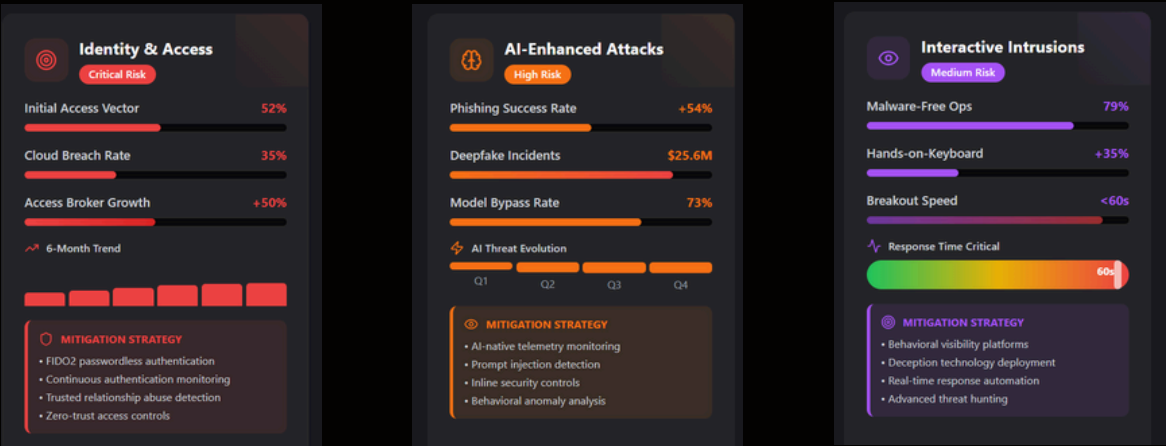
AI Threats & Market Relevancy Dashboard

Threat Intelligence Feed



Critical Attack Vectors & Mitigations

Real-time threat landscape analysis with mitigation strategies





2

AI has crossed the threshold from experimental to mission-critical — and attackers noticed first.

The enterprise adoption curve is steep, with over 20 major large language models (LLMs) launched in the past 18 months (e.g., *GPT-5*, *Claude 3.7*, *Gemini 2.5 Pro*, *Llama 4 Scout*, *Qwen 3*, *Grok 4*). Each introduces a new security and compliance surface.

- **AI-powered phishing works better:** Studies show LLM-generated phishing emails achieve 54% click-through, compared to just 12% for human-crafted lures.
- **Cloud AI abuse is real:** Attackers now resell compromised AI API keys in underground markets, a practice known as LLMjacking, fueling unauthorized large-scale queries.
- **AI as fraud enabler:** Deepfake voices and videos enabled fraud schemes that redirected up to \$25.6M in a single incident.
- **AI retraining risks:** Open-source LLMs like Qwen 2.5 were retrained to bypass Microsoft Defender with >8% success rates.
- **Insiders with AI augmentation:** Adversaries coached job candidates with AI to infiltrate enterprises — and once inside, used AI tools to blend in and exfiltrate data.

CISO Implication: Enterprises cannot treat AI adoption as “just another IT deployment.”

AI systems are both assets to protect and attack vectors to defend against. Governance, runtime controls, and observability are no longer optional — they are foundational.

3. Generative AI in Adversary Operations

AI Threat Evolution

- **LLMs generated phishing emails** with 54% higher click-through rates than human-crafted ones.
- **Deepfake voices and videos** powered multi-million-dollar fraud attempts (one case diverted \$25.6M).
- **Open-source models** (e.g., Qwen 2.5) were retrained to bypass endpoint security.

Implication: AI is lowering barriers to sophistication — the “script kiddie” of yesterday can now deploy nation-state-grade influence campaigns.

Nation-State Activity Indicators

- **China-linked intrusion activity rose** 150% overall, and 200–300% in finance, manufacturing, and media.
- **DPRK’s FAMOUS CHOLLIMA** leveraged AI-coached insider workers in 40% of cases.
- **Russia- and Iran-aligned operators** ran coordinated disinformation campaigns using LLM-generated content.

Implication: Adversary nations are scaling AI like enterprises — CISOs must expect long-term, persistent campaigns, not opportunistic hits.

Key Insight: Adversary nations are scaling AI like enterprises. CISOs must expect long-term, persistent campaigns, not opportunistic hits.

4. Where CISOs Must Focus in 2025

- **Identity & Access:** Defend against access-broker ecosystems with continuous authentication, password less methods (FIDO2), and monitoring of trusted relationship abuse.
- **Sub-Minute Response:** Breakout times of under a minute demand inline controls at the AI proxy layer, not just EDRs.
- **AI Governance:** With 20+ enterprise LLMs now in circulation, governance is no longer optional — CISOs must map which models are authorized for which business tasks.
- **Insider Risk Management:** Screening must evolve to detect AI-assisted deception in recruitment; technical controls (e.g., proxy enforcement) become key.
- **SOC Modernization:** Logs must include AI-native telemetry — prompt injection attempts, adversarial probe detection, and agent manipulation.



5. Case Studies in AI Attack Vectors

Case: Cloud Credential Abuse → AI Hijacking

- Attackers stole API keys from a consulting firm's cloud tenant, then resold "AI queries-as-a-service" in criminal forums.
- **Lesson:** Treat AI API keys like privileged credentials — rotate, vault, and monitor.

Case: Calendar Injection via AI Prompt

- A prompt injection campaign embedded in calendar invites hijacked smart-home controls through a connected AI assistant.
- **Lesson:** Prompt inputs are the new supply chain risk — enforce inline sanitization.

Case: Insider Developer with AI Assistance

- Adversaries placed AI-coached "employees" into victim firms, feeding stolen code into external AI systems for exfiltration.
- **Lesson:** Background checks aren't enough — AI activity monitoring must be integrated into DevSecOps.

6. The Aldome™ Lens

- **Analogy:** Aldome is to AI what Okta is to identity and Datadog is to observability.
- **Product Highlights:**
 - **SDK + Proxy dual mode** → deep CI/CD integration.
 - **SOC-ready deception** — SentinelTrap™.
 - **Air-gapped OS SecOps deployment** — Full Admin infrastructure.
 - **Not just SaaS filtering** — but agent and LLM model observability.
 - **Edge Security** — Mobile (BYD) enterprise chat application
- **Strategic Benefit:** CISOs can align AI adoption with regulation (GDPR, HIPAA, EU AI Act) while actively reducing breach risk.

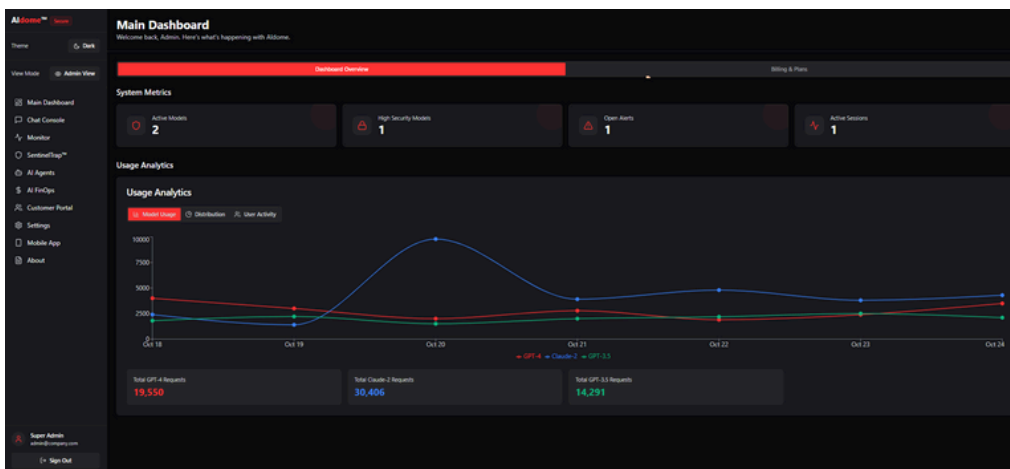


Image: Aldome Admin

7. Recommendations for CISOs

Immediate (0–3 months)

- Run an AI security workshop to map existing model usage.
- Audit access brokers / account hygiene.
- Deploy inline monitoring of AI prompts.

Medium Term (3–12 months)

- Extend SOC with AI-native logs.
- Establish central AI governance policy.
- Pilot deception honeypots targeting adversarial AI probes.

Strategic (12–24 months)

- Evolve AI risk posture into enterprise risk governance.
- Integrate AI-native security OS (like Aldome™) as part of core SecOps stack.
- Prepare for compliance scrutiny (AI Act, ISO AI).



8 Final Word

- **2025 is the year adversaries became enterprises** — monetizing access, embedding insiders, and weaponizing AI. CISOs must treat AI as critical infrastructure.
- **Aldome™ enables defenders to keep pace, embedding AI-native controls where they matter most:** at the model, the agent, and the prompt.

Aldome™ Strategic Defense Matrix

Comprehensive AI security posture aligned with 2025 threat landscape

IMMEDIATE RESPONSE

- Sub-minute breakout detection
- Real-time prompt monitoring
- Inline policy enforcement

AI GOVERNANCE

- Model authorization mapping
- Compliance automation
- Risk scoring engine

THREAT INTELLIGENCE

- Nation-state tracking
- Insider threat detection
- Deception technology

CONTACT

 INFO@AIDOME.CO

Disclaimer

This report is provided for informational purposes only and does not constitute legal, regulatory, investment, or professional advice. The threat intelligence data and statistics referenced herein are derived from third-party industry sources and publicly available information believed to be reliable at the time of publication. Aldome™ make no representation or warranty, express or implied, regarding the accuracy, completeness, or fitness for a particular purpose of the contents of this document. Security threats, vulnerabilities, and adversary tactics evolve rapidly; therefore, the findings and recommendations in this report should not be considered exhaustive or definitive. Organizations should perform their own independent assessments and consult qualified security professionals when making decisions related to AI, cybersecurity, and regulatory compliance. Aldome™ disclaim any liability for actions taken or not taken based on the information provided in this report.

