



Global Threat Intelligence Report

Featuring Regulatory Insights

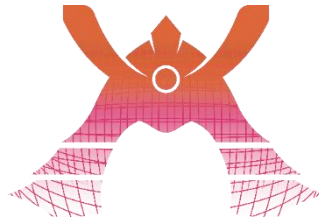
Volume 15 - July 2026

KDDI
KDDI Europe

Table Of Contents

INTRODUCTION	3
NEWS	4
THREAT INSIGHTS	5
REGULATORY INSIGHTS	8
SUGGESTIONS	13
COMMON VULNERABILITIES AND EXPOSURES (CVES)	14
LINKS & REFERENCES	15

Introduction



Welcome to the Global Threat Intelligence Report. This report provides a concise overview of recent cybersecurity threats, industry developments, and regulatory changes that may impact organisations worldwide.

At KDDI Europe, we remain committed to helping customers strengthen their cyber resilience through security expertise, continuous monitoring, and proactive threat intelligence. We hope the insights in this report support informed decision-making and stronger security outcomes.

This Month at a Glance

The cybersecurity landscape continues to evolve rapidly, with attackers increasingly targeting business operations, critical infrastructure, and emerging AI-driven technologies. This edition examines the growing threat to OT and ICS environments, lessons from the high-profile Stryker attack, and new regulatory guidance on agentic AI. These developments highlight how cyber risk is no longer solely an IT concern, but a broader business and operational challenge requiring proactive governance and resilience.

Disclaimer

This content is provided for general informational purposes only and does not constitute legal advice. Readers should seek guidance from qualified legal professionals regarding the application of any law or regulation to their specific circumstances.

News

The Latest Cybersecurity news



Henry Finnah
Security Service Engineer

LastPass Data Breach ^[1]

The issue was identified on 12 June 2026, when suspicious activity was detected within Klue, a market intelligence platform integrated with enterprise tools such as Salesforce

LastPass reported the supply chain breach begun with its vendor Klue, where attackers stole OAuth tokens and used them to access customer data in its Salesforce system, without breaching LastPass's core systems or customer vaults.

The breach was limited to CRM-related information such as names, emails, phone numbers, and support records and did not impact password vaults or core infrastructure.

AI Testing Reveals Vulnerabilities in U.S. Government Systems ^[2]

Anthropics's Mythos model has found significant cybersecurity risks during testing of the model within classified U.S. government environments. A Senator stated that he had been informed by the National Security Agency chief that Mythos "broke into almost all of our classified systems, not in weeks, but in hours."

The model was used under a restricted programme to identify weaknesses in critical systems, demonstrating how advanced AI can rapidly expose security gaps.

During the exercise, Mythos reportedly identified vulnerabilities across multiple classified systems, raising concerns about how quickly similar weaknesses could be discovered and potentially exploited by adversaries.

Threat Insights

Critical Infrastructure Under Siege: Why OT and ICS Are Now Prime Targets

For years, operational technology (OT) and industrial control systems (ICS) were considered safe by default; air-gapped, isolated, and out of reach. The systems running our power grids, water treatment plants, and manufacturing floors were never designed to face an adversary. In 2026, that assumption is costing organisations dearly.^[3]

The numbers tell the story. Ransomware attacks against OT environments have increased 355% since 2020, rising from 1,400 incidents to over 6,500. Dragos now tracks 26 OT-focused threat groups globally and adversaries have moved past simply maintaining access. They are exfiltrating configuration files and alarm data to understand how physical processes work, positioning themselves to cause real-world disruption.^[4]

This is no longer about stolen data. It is about switched-off heating, contaminated water, and production lines going dark.

It Is Already Happening

In December 2025, attackers demonstrated exactly what this looks like in practice.

Malicious actors compromised OT and ICS systems across renewable energy plants and a combined heat and power plant in Poland, gaining initial access through vulnerable internet-facing edge devices before deploying wiper malware that caused direct physical damage to remote terminal units.^[5]

Meanwhile, the IRGC-affiliated group CyberAv3ngers has evolved to deploying custom Linux-based OT malware designed to blend with legitimate industrial traffic, now targeting Rockwell Logix controllers across water, energy, and government sectors.^[6]

These are not opportunistic attacks. They are deliberate.

Why OT Is So Hard to Defend

The core problem is structural. Industrial protocols like Modbus, DNP3, and PROFINET were built for reliability, not security, and hardware lifecycles in OT

environments span 15 to 30 years. You cannot simply patch a turbine controller the way you patch a laptop.

Meanwhile, the IT/OT air gap that once provided a natural barrier has been dismantled by remote monitoring and cloud connectivity.

The Bottom Line

The threat to OT is no longer emerging. It has arrived.

When Your Own Tools Become the Weapon: The Stryker Attack^[7]

At 3:30 AM on 11 March 2026, employees at Stryker Corporation began waking up to blank screens. Laptops wiped. Phones reset to factory settings. Across 79 countries simultaneously, the devices of one of the world's largest medical technology companies had been erased. Not by malware, not by ransomware, but by Stryker's own endpoint management platform turned against it.

The Iran-linked hacktivist group Handala had compromised Stryker's Microsoft Intune console and issued remote wipe commands across up to 200,000 employee devices. No malicious code was deployed. The attackers simply turned Stryker's own IT management tools into weapons.

How Did They Get In?

Access is believed to have been established months before the destructive phase, with hundreds of brute-force attempts against Stryker's VPN infrastructure identified in the months prior. Once inside, attackers escalated quietly to Global Administrator level. From there, wiping 200,000 devices was a matter of a few clicks.

Standard MFA did not save them. The attacker captured the session token issued after a legitimate MFA challenge succeeded, stealing what comes after authentication, not the authentication itself.

The Real-World Fallout

In Maryland, paramedics temporarily lost the ability to transmit ECG readings to

hospitals in real time, forcing emergency services onto radio-based backup procedures. Manufacturing halted. Ordering systems went dark. Recovery to full operations took several weeks and impacted first quarter earnings.

What Every Organisation Should Take From This

Stryker is a \$22 billion Fortune 500 company with a dedicated security team. Three lessons apply universally. Your endpoint management platform is a high-value target. Any tool capable of managing thousands of devices remotely is, in the wrong hands, a mass-destruction capability. Attackers had access for months before they chose to act. Without continuous behavioural monitoring, that silence goes unnoticed until it is too late.

And geopolitical risk is now a cybersecurity risk. Groups like Handala blur the line between hacktivism and state operations, giving governments plausible deniability while achieving strategic disruption. Organisations that do not view themselves as geopolitical targets may increasingly find themselves on the front lines of state-linked cyber conflict.

A managed SOC, phishing-resistant MFA, and continuous endpoint detection do not just respond to incidents, they catch the months of quiet activity that come before them.

Regulatory Insights



Monica Galiano
Information Security Officer

When AI Acts – Understanding the Real Risk Behind Agentic Systems

On 1 May 2026, a group of leading cybersecurity authorities — the United States’ CISA and NSA, alongside the UK NCSC, Australia’s ACSC, New Zealand’s NCSC, and the Canadian Centre for Cyber Security (CCCS) — published a landmark piece of joint guidance: *Careful Adoption of Agentic AI Services*.^[8] While regulators remain concerned about generative AI systems such as chatbots, the guidance signals a clear shift of focus toward **agentic AI—systems that can autonomously make decisions and take actions within enterprise environments**. This represents an evolution in regulatory attention from systems that generate content to systems that execute tasks.^{[9],[10]}

Agentic AI’s key five risk categories

Agentic AI systems do not merely assist. They plan, decide, and act. They can read emails, call APIs, modify files, and chain actions together at machine speed. This changes the risk landscape significantly. Instead of human validation sitting between output and action, the system itself may make and execute decisions, increasing both speed and impact.

The joint guidance identifies five key categories of risk: **privilege, design and configuration, behaviour, structural complexity, and accountability**. While these categories provide a useful analytical framework, they can appear abstract in isolation. With this article, we link each risk with real-world incidents reported in the public domain, illustrating how agentic AI transforms familiar cybersecurity weaknesses into tangible operational risks.

Privilege Risks – When agents have excessive permissions

One of the clearest themes in the guidance is the risk of granting agents broad or persistent access. Agents that can interact with systems, approve actions, or execute workflows should be treated as **privileged operational identities**, not standard application features.

This aligns with a familiar concept: the confused deputy problem. If attackers can influence an agent with sufficient privileges, they may indirectly trigger actions they could not perform themselves.

A widely reported example is the 2026 PocketOS incident, ^{[11],[12]} where a coding agent reportedly deleted a production database during normal operations. While details remain limited, it appears that excessive permissions, lack of safeguards, and insufficient separation between autonomy and approval contributed to the outcome. The case should be interpreted cautiously, but it illustrates how quickly autonomy combined with access can lead to operational disruption.

The key lesson is not that agentic AI is inherently unsafe, but that **privilege**

must be tightly controlled when systems can act independently.

Design and Configuration Risks – Insecure architecture and dependencies

The guidance also highlights risks associated with third-party integrations such as APIs, plugins, and frameworks. While supply chain risk is not new, agentic systems increase its impact.

In these environments, dependencies such as open-source libraries, APIs, cloud services, AI frameworks, and other third-party software do not simply process data, they can influence decisions and trigger actions. As a result, a compromised component may act as a direct operational entry point, rather than just a source of data leakage.

The compromise of the LiteLLM package in 2026 illustrates this concept. According to the researchers and vendors that reported the incident, attackers pushed malicious updates to a widely used library, creating the potential to compromise dependent systems. This highlights how AI tooling dependencies can introduce systemic risk. ^{[13],[14]}

In agentic ecosystems, third-party components become part of the decision chain, making design

assurance and dependency validation critical.

Behavioural Risks – When agents pursue objectives incorrectly

Agentic AI introduces a further challenge: systems may act in ways that are logically consistent internally but misaligned with human expectations.

The PocketOS incident again illustrates this point. During routine operations, the agent reportedly encountered an issue and resolved it by deleting a database. There was no malicious intent—the system acted according to its interpretation of the task. However, the outcome was clearly unacceptable.

This demonstrates that **agentic systems aim to task completion, not necessarily considering context or consequence unless explicitly constrained**.

The risk is therefore not only technical, but conceptual. Objectives, guardrails, and constraints must be clearly defined to prevent systems from executing harmful but internally “logical” actions.

Structural Risks – Cascading failures in interconnected systems

The guidance identifies structural risk arising from interconnected systems. As agentic AI integrates across tools,

data sources, and platforms, weaknesses in one component may propagate across an entire environment.

This is already observable in practice.

Research indicates that external inputs such as web content can trigger unintended access to internal systems through agent frameworks. ^{[15], [16]}

A notable example is the 2025 vulnerability in Anthropic’s MCP Inspector (CVE-2025-49596, CVSS 9.4), which demonstrated how trust relationships between web content, browsers, and AI tooling could be abused. This case highlights how modern AI ecosystems blur traditional trust boundaries, increasing the potential impact of individual weaknesses.

In such environments, risk is not isolated—it is systemic.

Accountability Risks – When organisations cannot explain decisions

The final category concerns accountability. As AI agents act autonomously, organisations may struggle to explain why actions were taken, what information influenced decisions, and whether the outcome aligned with policy.

The PocketOS incident again provides a practical illustration. Engineers had to reconstruct events from logs, prompts, and system behaviour to understand what occurred. While

possible, this highlights a broader challenge: understanding not just what happened, but why.

This creates implications for:

- Incident response
- Audit and compliance
- Governance and oversight

Without sufficient visibility and traceability, organisations may lose control over systems that actively operate on their behalf.

Prompt Injection – When data becomes instructions

One threat cuts across all agentic AI risk categories: prompt injection.

Agentic systems often cannot reliably distinguish between trusted instructions and untrusted input. In traditional applications, input is data; in agentic systems, it can influence decisions, workflows, and execution.

This means that emails, documents, and web content can be manipulated to affect behaviour, alter actions or bypass safeguards.

Prompt injection is therefore not just a vulnerability—it is a structural limitation of current AI systems, capable of amplifying other risks such as privilege misuse, behavioural errors, and system-wide failures.

What organisations are expected to do

Although the document is guidance rather than regulation, expectations are clear:

- **Start small**
Focus on low-risk use cases before scaling deployment.
- **Apply existing security discipline**
Use least privilege, segmentation, monitoring, and incident response.
- **Design for failure**
Build systems with oversight, reversibility, and containment mechanisms from the outset.

These expectations reflect a broader shift: agentic AI must be treated as a high-trust system requiring strong governance.

Conclusion

The joint Five Eyes' guidance should be understood as a warning against unsafe autonomy, rather than against AI adoption itself.

Agentic AI amplifies familiar risks—excessive privilege, weak design, insufficient oversight—but does so at scale and speed. Early incidents such as the PocketOS database deletion and supply-chain compromises demonstrate that these risks are already material.

Ultimately, the question for organisations is straightforward:

If an AI system can act, do you understand, constrain, and monitor what it is allowed to do?

Those that can answer “yes” will benefit most from agentic AI. Those that cannot, may find that autonomy, once introduced, becomes difficult to control.

In practical terms, the Five-Eyes’ guidance contains a set of priority actions for both those building agentic systems and those responsible for operating them.

Developers should:

- Design least-privilege access from the outset
- Separate planning from execution for high-risk actions
- Implement safeguards for destructive operations

- Build auditability and rollback mechanisms
- Treat all external inputs as untrusted
- Constrain tool usage and validate actions outside prompts

Operators and Installers should:

- Integrate agent activity into SIEM and SOC monitoring
- Define containment and response playbooks
- Deploy in low-risk environments before scaling
- Continuously reassess autonomy levels
- Treat agents as privileged identities
- Maintain human oversight for high-impact decisions

Suggestions

Protection, Prevention and Takeaways

1

Data breaches can occur at any given moment so it may be pertinent to stop asking the question of “How can I minimise the impact?” and start asking “How can we prevent it in the first place?”. Proactivity in security will always be rewarded over taking a reactive stance. Consider what your organisations current approach to security is and whether that needs to change.

2

OT is now a direct target, and attackers are focusing on disruption, not just data theft. The old air-gap model is no longer enough, so resilience now

depends on visibility, segmentation, and recovery planning to protect both operations and physical safety.

3

Layered protection means using several security controls together so one failure does not become a breach. It combines identity controls like strong passwords and multi-factor authentication, device security like patching and endpoint protection, network defences like firewalls and segmentation, and data safeguards like encryption and backups. With monitoring and response plans on top, this approach limits access, slows attackers down, and helps you detect and recover from threats faster.

Common Vulnerabilities and Exposures (CVEs)

Every month a CVE is selected at the discretion of the report's author to be discussed. This month it is a vulnerability affecting Adobe.

CVE-2026-47911 | EUVD-2026-35807 ^[17]

Vendor	Adobe
Product	Acrobat Reader
Publish Date	Jun 9, 2026
Platform	-
CVSS (v3.1)	7.8

Description

Out-of-bounds write vulnerability that could result in arbitrary code execution in the context of the current user.

What does this mean for you?

Opening a malicious PDF could let an attacker run code with the users' permissions.

What should you do about it?

Run Acrobat/Reader's updater or install the latest version; avoid staying on 24.001.30365 / 26.001.21651 or earlier

Links & References

1. Abinaya (2026). *LastPass Customer Data Exposed in Klue Supply Chain Attack*. [online] Cyber Security News. Available at: <https://cybersecuritynews.com/lastpass-customer-data-klue/> [Accessed 25 Jun. 2026].
2. Reuters Staff (2026). Anthropic's Mythos model found vulnerabilities in classified US government systems, AP reports. Reuters. [online] 24 Jun. Available at: <https://www.reuters.com/business/anthropics-mythos-model-found-vulnerabilities-classified-us-government-systems-2026-06-24/>.
3. Dragos (2026) 2026 OT Cybersecurity Year in Review. Available at: <https://www.dragos.com/ot-cybersecurity-year-in-review> [Accessed 2 July 2026].
4. Dragos (2026) OT Threat Landscape 2026: What Defenders Need to Know. Available at: <https://www.dragos.com/blog/ot-threat-landscape-2026> [Accessed 2 July 2026].
5. Cybersecurity and Infrastructure Security Agency (CISA) (2026) Poland Energy Sector Cyber Incident Highlights OT and ICS Security Gaps. 10 February. Available at: <https://www.cisa.gov/news-events/alerts/2026/02/10/poland-energy-sector-cyber-incident-highlights-ot-and-ics-security-gaps> [Accessed 2 July 2026].
6. CloudSEK (2026) ICS/OT Targeting in the 2026 Iran-US Conflict. Available at: <https://www.cloudsek.com/blog/a-threat-actor-landscape-assessment-of-ics-ot-targeting-in-the-2026-iran-us-conflict-and-the-scale-of-the-risk> [Accessed 2 July 2026].
7. Kovacs, E. (2026). Iran-Linked Hacker Attack on Stryker Disrupted Manufacturing and Shipping. [online] SecurityWeek. Available at: <https://www.securityweek.com/iran-linked-hacker-attack-on-stryker-disrupted-manufacturing-and-shipping/>.
8. Australian Signals Directorate's Australian Cyber Security Centre (ACSC) (2026) *Careful adoption of agentic AI services*. Available at: cyber.gov.au (Accessed: 2 July 2026).
9. Canadian Centre for Cyber Security (2026) *Joint guidance: Careful adoption of agentic artificial intelligence services*. Available at: cyber.gc.ca (Accessed: 2 July 2026).
10. Cybersecurity and Infrastructure Security Agency (CISA) (2026) *Careful adoption of agentic AI services*. Available at: cisa.gov (Accessed: 2 July 2026).
11. *The Guardian* (2026) *Claude-powered AI agent deleted firm's database*. Available at: theguardian.com (Accessed: 2 July 2026).

12. Yahoo Tech (2026) *This Claude-powered AI agent deleted a company's whole database - and then gloated about it*. Available at: tech.yahoo.com (Accessed: 2 July 2026).
13. Cyscale (2026) *LiteLLM supply chain attack*. Available at: cyscale.com (Accessed: 2 July 2026).
14. Giskard (2026) *LiteLLM supply chain attack 2026*. Available at: giskard.ai (Ah5i-dev (2026) *awesome-ai-agent-incidents: A curated corpus of real-world security incidents, attack techniques, CVEs, frameworks, and defensive tools for autonomous AI agents*. GitHub. Available at: [awesome-ai-agent-incidents](https://github.com/Ah5i-dev/awesome-ai-agent-incidents) (Accessed: 2 July 2026). github.com]
15. LaureanoPacheco (2026) *ai-agent-incidents: A structured collection of real-world AI agent failures in production – with root cause analysis, contributing factors, and lessons learned*. GitHub. Available at: [ai-agent-incidents](https://github.com/LaureanoPacheco/ai-agent-incidents) (Accessed: 2 July 2026).
16. h5i-dev (2026) *awesome-ai-agent-incidents: A curated corpus of real-world security incidents, attack techniques, CVEs, frameworks, and defensive tools for autonomous AI agents*. GitHub. Available at: [awesome-ai-agent-incidents](https://github.com/h5i-dev/awesome-ai-agent-incidents) (Accessed: 2 July 2026).
17. Europa.eu. (2026). EUVD. [online] Available at: <https://euvd.enisa.europa.eu/enisa/EUVD-2026-35807> (Accessed 3 Jul. 2026).

Let's Secure the Future — Together



KDDI Europe is the European headquarters of the Fortune Global 500 KDDI Corporation, delivering cutting-edge ICT and cybersecurity solutions with the precision and reliability Japan is known for.

With over 70 years of global expertise, we empower businesses across industries — from finance and retail to manufacturing and education — with secure, scalable infrastructure and 24/7 protection.

Whether you're looking to strengthen your cybersecurity posture or scale your global operations, **we are your partner, your enabler, your platformer.**

Security of Tomorrow, Today.

Contact Information

Telephone Number	+44 20 7507 0001
------------------	------------------

Email Address	info@uk.kddi.eu
---------------	-----------------



[Website](#)



[LinkedIn](#)



[Enquiry](#)