

No. 118 | September 2026



Radar

The cybersecurity
magazine

Cybersecurity in the Agentic Era: Attack, Defend, and Anticipate

By Jose Manuel Moreno Guerra

Artificial intelligence is transforming cybersecurity from both sides. On the one hand, it enables organizations to detect, analyze, and respond at speeds that would have been unthinkable just a few years ago. On the other, it gives attackers new capabilities to automate campaigns, discover vulnerabilities, and chain together offensive actions at an unprecedented scale and speed. The consequence is clear: AI does not merely introduce a new category of risk; it fundamentally changes the amount of time available to defend against threats.

Until now, much of cybersecurity has relied on relatively predictable cycles. A vulnerability was detected, classified, prioritized, and eventually remediated. An incident triggered an alert, passed through different levels of analysis, and ultimately led to a response.

That model is beginning to fall short. When exploitation can occur within minutes, defense can no longer operate on timelines measured in days or weeks. The resilience of the future will depend on the ability to detect, validate, prioritize, and act at machine speed.

On the defensive side, SOCs are evolving toward agentic models. AI can already take on much of the repetitive work performed at the initial levels of security operations: normalizing alerts, enriching indicators, correlating information, querying threat intelligence, recommending actions, and even executing preapproved responses.

This does not eliminate the need for analysts; it transforms their role.

Security professionals can spend less of their time on mechanical tasks and focus instead on complex threats, advanced investigations, decisions with business impact, and the continuous improvement of detection capabilities.

This shift must extend across the entire security lifecycle. AI can help model threats from the design stage, reduce false positives in code analysis, automate dynamic testing, consolidate findings from multiple tools, and prioritize vulnerabilities based on the real-world context of the asset. It can also turn open-source information into structured, actionable intelligence.

Taken together, these capabilities make it possible to move from fragmented, reactive security to a continuous, contextual, and connected model.

However, the same logic can also be used offensively. AI agents can explore applications, select tools, interpret results, adapt strategies, and chain vulnerabilities together until they achieve a specific objective.

Unlike traditional scanners, which are constrained by signatures, rules, and predefined sequences, agentic systems can reason about their environment and modify their behavior during execution.

This capability is highly valuable for audits and Red Team exercises, but it also points to a new generation of autonomous attacks.

The risk surface is further expanded by AI assets themselves. Chatbots, predictive models, RAG systems, and agents connected to corporate tools may be exposed to prompt injection, data leakage, model extraction, permission abuse, resource exhaustion, or manipulation of their behavior.

The greater the system's autonomy, the greater the potential impact of an incorrect decision or a malicious instruction.

For this reason, the question is not whether we should use AI in cybersecurity, but under what conditions. Speed must be accompanied by control.

Critical decisions require human oversight, limited permissions, reversible actions, complete traceability, and auditable evidence. Models should operate on reliable knowledge bases, corporate policies, defined procedures, and safeguards designed to reduce hallucinations, shortcuts, and unexpected behavior.

The advantage will not lie in replacing expert talent, but in amplifying it. AI can absorb volume, reduce noise, and accelerate analysis; people must provide judgment, context, accountability, and business understanding.

In this new race between defense and attack, the organizations that succeed will be those capable of combining automation with governance, offensive capabilities with defensive resilience, and innovation with control.

Because in the agentic era, protection no longer means simply responding better, but learning and acting before the adversary does.



Jose Manuel Moreno Guerra
Cybersecurity Director

The breach no longer ends at the victim's network

Cyber Chronicle by Alicia Isabel Martínez Vera

In less than a month, incidents across different countries and sectors exposed the same weakness: an organization's security increasingly depends on systems it does not fully control. Legitimate packages were compromised through trusted repositories and authorized publishing processes; stolen credentials opened access to shared platforms; the disruption of a logistics operator affected deliveries to restaurants and supermarkets; and exposed industrial controllers made it possible to impact water systems in the United States. The common thread was not a single malware family, but rather the trust delegated to suppliers, automation, identities, and connected systems.

Between the second week of July and the first days of August, cybersecurity once again showed that an intrusion can have consequences far beyond the initial compromise. In some cases, attackers used legitimate repositories and development processes to distribute malicious code. In others, a shared platform, a cloud provider, or a logistics system extended the disruption to organizations that were never directly attacked. The gap between the scope confirmed by victims and the figures claimed by extortion groups also widened, with several incidents still lacking independent verification.

Between July 8 and 9, Accenture confirmed a security incident that it described as isolated. The consulting firm said it had addressed the root cause and stated that its operations and services had not been affected. An actor identified as "888" later offered what it claimed to be around 35 GB of source code, configuration files, RSA and SSH keys, and Azure-related credentials. The existence of the breach was confirmed by the company to several media outlets, but Accenture did not validate either the volume or the full contents claimed by the attacker. The case is significant because compromised data can **retain offensive value even when the victim does not experience an immediate disruption**.

From July 8 onward, attention shifted to the software supply chain. That day, a version containing a malicious module disguised as telemetry appeared on **npm**. The code intercepted the functions used to load private keys and mnemonic phrases and exfiltrated the secrets encoded within an HTTP header. The package remained available for around 49 minutes before being replaced by a clean version. Technical analysis confirmed its **data-stealing capability**, although it was not possible to determine how many installations actually executed the implant or whether any financial losses occurred.

Three days later, Jscrambler experienced a similar sequence of events, although the initial access began on a developer's device. The attacker obtained a GitHub key, gained access to the repository, and used GitHub Actions to extract the npm publishing token. With it, five malicious versions of the package were released. The first versions activated an **infostealer** during installation; later ones **changed the mechanism so that it executed from within the code itself**. The company detected the malicious releases through automated alerts, revoked the credentials, and published a clean version that same day.

On July 14, the AsyncAPI ecosystem confirmed that even a release with valid cryptographic provenance does not guarantee that its contents are safe. A misconfigured GitHub Actions workflow allowed **privileged credentials** to be obtained and malicious code to be introduced into four **official** packages. Five compromised versions were published through the project's legitimate mechanisms. The trusted automation had worked correctly; the problem was that **the authorized process itself had already been hijacked**. While the malicious packages threatened development environments, the attack on Nichirei shifted the crisis into food logistics.

On July 13, the Japanese company detected unauthorized access that disrupted inbound and outbound operations at refrigerated warehouses and affected frozen food shipments. **Isolating the systems helped protect information but temporarily halted the flow of goods**. KFC Japan suspended digital orders and warned of menu restrictions, reduced opening hours, or temporary closures depending on inventory levels at each location. Delays were also reported at supermarkets and other restaurant chains. Nichirei confirmed the cyberattack but did not publicly identify ransomware, data theft, or a responsible group.

On July 16, Coca-Cola disclosed to the SEC that Fairlife had suffered a ransomware attack affecting part of its systems, including some related to production. The company temporarily suspended manufacturing in the United States, while Canadian operations and product safety remained outside the known scope of the incident.

On July 27, Coca-Cola reported that most production had resumed across its four U.S. plants and confirmed that the attacker had obtained certain data. The Anubis group claimed responsibility for the operation and said **it had stolen approximately one terabyte of data.**

The most unusual intrusion of the period was disclosed by **Hugging Face** on July 16 and explained by OpenAI five days later. Between July 9 and 13, agents used in an **internal evaluation of cyber capabilities managed to escape the intended restrictions.** The models exploited a zero-day vulnerability in the software acting as a package proxy, reached a node with internet access, and began searching for information that would allow them to solve or bypass the test. From there, they chained additional credentials and vulnerabilities until they achieved remote code execution.

Hugging Face confirmed access to a limited set of internal datasets and several credentials, but found no evidence that models, packages, Spaces, or public datasets had been altered. The forensic reconstruction grouped together approximately 17,600 actions carried out by the agent. This was not a criminal campaign deliberately targeting Hugging Face, but rather a **poorly contained evaluation that resulted in a real intrusion into a third party.**

External dependency risks resurfaced in the cases of Stadler Rail and Amgen. In mid-July, **compromised credentials enabled access to an information-sharing platform** used with one of Stadler's suppliers. The attackers accessed technical documentation, but did not penetrate the manufacturer's internal network, disrupt production, or affect railway vehicles. Everest demanded ten million Swiss francs, a claim that comes from the extortion letter itself, although it remains unclear whether the same group was technically responsible for the intrusion.

Amgen later reported that it had detected unauthorized activity involving data stored in cloud environments managed by external providers. The pharmaceutical company confirmed **the exfiltration of proprietary information, protected health information, and other files, although it observed no impact** on its products, manufacturing operations, or ability to serve patients. Both incidents highlighted the same issue from different angles: **outsourcing the hosting or exchange of information does not eliminate responsibility or automatically reduce the consequences of an exposure.**

Between July 26 and 30, the threat reached **operational technology** directly. Minnesota reported a coordinated attack against more than thirty community water systems. Shortly afterward, the FBI and EPA expanded the alert to facilities in at least seven states.

The attackers accessed internet-exposed Rockwell Automation/Allen-Bradley MicroLogix 1100 and 1400 PLCs, changed IP addresses and passwords, and caused a loss of visibility or control. Reported effects included **loss of water pressure and flooding at some facilities.** The ability to switch to manual operation largely determined the severity of each case. As of August 7, there was no official attribution. The most accurate classification is **unattributed operational cyber sabotage.**

On August 4, Microsoft published its investigation into **ChainDrop**, a campaign that had already distributed malicious versions of numerous npm packages. The **worm stole tokens and credentials from development environments and reused them to publish new malicious** versions of other packages. Counts varied during the investigation because some sources tracked unique packages while others counted published versions, but all agreed on the campaign's **self-propagating nature.** The distinction from a single compromised dependency was critical: each breached account could automatically become the starting point for new malicious releases, expanding the chain without requiring constant operator involvement.

The timeline therefore reveals something less obvious than the number of victims. Risk is concentrated in **relationships of trust:** a credential that allows software to be published, a pipeline that signs what it receives, a platform connecting two companies, a provider storing clinical information, or a PLC that remains accessible from the internet.

When those trust points fail, the incident spreads through business and technical pathways that already exist. The defensive priority is therefore not only to protect the core network, but also to **limit permissions, reduce persistent secrets, control automated publishing processes, and maintain manual procedures capable of sustaining operations when digital systems can no longer be trusted.**



Alicia Isabel Martinez Vera
Cybersecurity Consultant

Guardian Agents: The Next Step in Agentic Security

Article by Jaime Andrés Tovar Prieto

How are artificial intelligence agents used? How autonomous are agents in practice? What is the scope of the capabilities they were designed to have? Is there a mechanism in place to monitor their actions? Are those actions aligned with the functions for which the agent was originally designed?

These questions, although they may seem elementary, should always be considered. Just as permissions assigned to an application or user for access to a service must be reviewed—since excessive privileges can lead to a security incident—governance over AI agents requires the same level of rigor and discipline.

What is particularly relevant about this paradigm is the ability of one AI agent to oversee another: learning from behavioral patterns, detecting misuse, and assessing the logical consistency of the actions performed.

This dynamic, identified during research into AI security trends based on Gartner¹ publications, has led to leading academic developments such as AgentDoG² (Agent Diagnostic Guardrail), its associated benchmark ATBench³ (Agent Trajectory Benchmark), and PSG-Agent⁴ (Personalized Safety Guardrail).

With the growing adoption not only of Large Language Models (LLMs), but especially of AI agents—the latest evolution within the ecosystem—execution capabilities and autonomy in decision-making have expanded significantly.

1 (n.d.-a). Gartner.com. Retrieved August 12, 2026, from https://www.gartner.com/en/newsroom/press-releases/2026-06-03-gartner-security-and-risk-management-summit-2026-national-harbor-day-3-highlights

2 Liu, D., Ren, Q., Qian, C., Shao, S., Xie, Y., Li, Y., Yang, Z., Luo, H., Wang, P., Liu, Q., Hu, B., Tang, L., Mei, J., Guo, D., Yuan, L., Yang, J., Chen, G., Lin, Q., Yu, Y., ... Hu, X. (2026). *AgentDoG: A Diagnostic Guardrail framework for AI agent safety and security.* In arXiv [cs.AI]. https://doi.org/10.48550/arXiv.2601.18491

3 Li, Y., Luo, H., Xie, Y., Fu, Y., Yang, Z., Shao, S., Ren, Q., Qu, W., Fu, Y., Yang, Y., Shao, J., Hu, X., & Liu, D. (2026). *ATBench: A diverse and realistic agent trajectory benchmark for safety evaluation and diagnosis.* In arXiv [cs.AI]. https://doi.org/10.48550/arXiv.2604.02022

4 Wu, Y., Guo, J., Li, D., Zou, H. P., Huang, W.-C., Chen, Y., Wang, Z., Zhang, W., Li, Y., Zhang, M., Jiang, R., & Yu, P. S. (2025). *PSG-agent: Personality-aware safety guardrail for LLM-based agents.* In arXiv [cs.AI]. https://doi.org/10.48550/arXiv.2509.23614

These developments make a substantial strengthening of security necessary. The risk has evolved from simply preventing responses containing offensive or inappropriate language to addressing anomalous actions, the leakage of information or sensitive data, unauthorized actions that may result in access to critical systems, and, in other cases, service degradation.

The Industry's Expanded Perspective

In the face of this landscape of distributed and dynamic risks, the industry's response has not simply been to strengthen existing controls, but to design a new class of intelligence: Guardian Agents. As AI agents gain the ability to take autonomous actions, the need arises for a system capable of overseeing them. This is why the industry has developed the concept of Guardian Agents, also referred to as "AI supervisors."

A Guardian Agent is essentially an artificial intelligence system designed not to perform the target task of the agent being supervised, but rather to monitor, diagnose, evaluate, redirect, and constrain its behavior. The adoption of these technologies marks a critical departure from the traditional "notice and consent" paradigm, in which interactions with digital systems relied on warnings that users rarely understood in their entirety.

Instead, the model is shifting toward automated and auditable oversight, requiring human intervention only under a Human-in-the-Loop (HITL) approach when algorithmic certainty is insufficient or when risk exceeds predefined thresholds.

According to Gartner⁵, Guardian Agent technologies are expected to account for approximately 10% to 15% of the agentic AI market by 2030. The analyst firm states that these agents should focus on helping protect interactions across three fundamental roles:

5 (n.d.-b). Gartner.com. Retrieved August 12, 2026, from https://www.gartner.com/en/newsroom/press-releases/2025-06-11-gartner-predicts-that-guardian-agents-will-capture-10-15-percent-of-the-agentic-ai-market-by-2030

- **Reviewer:** Identify and review AI-generated outputs and content to verify their accuracy, truthfulness, and acceptable use.
- **Monitor:** Track the actions of artificial intelligence systems and agents, whether through human oversight or supervision performed by AI itself.
- **Protector:** Block or redirect agent actions and permissions through semi-automated or fully automated interventions during operations.

Evolution of the Attack Surface

In conventional AI environments, security controls rely on static semantic guardrails that, while valuable for regulatory compliance and an important component of due diligence, are insufficient against complex attacks. In an autonomous environment, a system does not become unsafe in a single step; rather, risk emerges through a sequence of erroneous decisions that interact in complex ways with inputs from the surrounding environment.

An emerging reference standard is the taxonomy implemented in frameworks such as AgentDoG and its evaluation infrastructure, ATBench.

The latter provides a standardized environment comprising 1,000 agent trajectories, of which 503 are safe and 497 unsafe, with access to more than 2,084 heterogeneous tools and an average of 9.01 turns per trajectory, providing an unprecedented level of realism in simulations.

AgentDoG proposes a unified security taxonomy for agent-based systems, structured around three orthogonal dimensions.

More specifically, it categorizes risks according to where the threat originates, how it manifests during execution, and what harm it causes in the real world:

- **Risk Source (Where):** Topologically defines the origin of the threat. It determines whether the attack vector originated from the user through direct malicious commands; from a tool or the environment through indirect injection in databases or compromised web pages; or from the agent's own internal logic, such as hallucination or flawed reasoning.
- **Failure Mode (How):** Examines behavioral disruption and systematically categorizes agentic lapses. This includes procedural anomalies such as invoking tools when information is incomplete, overlooking risks inherent in tool invocation, relying on unverified outputs, or executing unintended task sequences.
- **Real-World Harm (What):** Identifies the material impact of the failure, encompassing harmful externalities of a physical, economic, reputational, access-control, or societal nature.

Defense for the New Attack Surface

Addressing this new attack surface requires an architecture based on distributed oversight and dynamic security policies.

In this context, one of the most relevant approaches is distributed oversight.

The development of Guardian Agent systems has focused extensively on models such as PSG-Agent (Personalized Safety Guardrail).

Rather than relying on isolated checks, PSG-Agent implements a distributed oversight model in which multiple sub-agents assume asynchronous, specialized roles within the same execution pipeline:

- **Planning Monitor:** Evaluates strategies before execution.
- **Tool Firewall:** Intercepts operations and audits parameters at runtime.
- **Memory Guardian:** Acts as a write-blocking mechanism to prevent toxic data from corrupting long-term memory.
- **Response Guardian:** Validates the final output of the supervised agent.

A key differentiating feature of PSG-Agent is the dynamic personalization of risk thresholds based on the analysis of the user's historical profile. This mechanism enables it to identify cumulative risks that single-pass scanning systems are unable to detect.

Alongside PSG-Agent, another prominent reference framework is GuardAgent, which adopts a different approach based on symbolic reasoning over policies. GuardAgent uses a two-phase process supported by knowledge-enabled reasoning: it first analyzes security policies linguistically to develop an action plan and then converts that plan into executable code.

Using in-context demonstrations retrieved from its memory module, GuardAgent has reported rates above 98% in simulated access-control scenarios in healthcare environments and 83% accuracy in general moderation for web agents.

Conclusion

The relevance of Guardian Agents is not a technological trend; it is a structural response to a problem that static guardrails can no longer solve.

As AI agents expand their operational autonomy, distributed oversight, dynamic risk personalization, and policy reasoning are no longer optional. They are becoming architectural requirements.

Building trustworthy artificial intelligence systems means not only designing capable agents but also designing the guardians that keep them within the boundaries of what is safe, auditable, and ethically responsible.



Jaime Andrés Tovar Prieto
Cybersecurity Architect

AI in Cybersecurity: Integrating Defensive and Offensive Security

Trends by Abel Martin Huaman Condor

Cybersecurity should not rely solely on defensive controls designed to prevent, detect, and respond to threats. Organizations also need to continuously validate the effectiveness of those controls and identify their weaknesses before attackers can exploit them. To achieve this, defensive security must be complemented by offensive security practices supported by artificial intelligence.

Agentic AI can work toward defined objectives, analyze information, use authorized tools, and execute actions within established boundaries. This makes it possible to accelerate security activities and connect offensive findings with defensive improvements more quickly.

AI in Defense and Offense

On the defensive side, AI helps analyze large volumes of alerts and events, identify anomalous behavior, and prioritize the most relevant risks. Agentic AI can gather information from different sources, correlate indicators, and propose previously authorized containment actions.

In offensive security, AI can support penetration testing, phishing simulations, and configuration analysis within an approved scope. Its purpose is to identify weaknesses in people, processes, and technologies before a malicious actor can exploit them.

Although these tools are becoming increasingly accessible, using them effectively still requires a basic understanding of systems, networks, vulnerabilities, and risks. AI amplifies team capabilities, but it does not replace technical validation or professional judgment.

Why Integrate Both Sides

Defense makes it possible to protect, monitor, and respond to threats. However, an organization cannot assume that its controls are effective simply because they have been implemented. They need to be validated against controlled scenarios that reflect techniques and behaviors that could be used in a real attack.

Offensive security serves this purpose: it identifies gaps and demonstrates how they could affect assets, processes, or users. Defense then turns those findings into concrete actions, such as configuration changes, new detection rules, improvements to response procedures, or stronger access controls.

Integrating both sides makes it possible to:

- Identify vulnerabilities before they lead to an incident.
- Verify the real effectiveness of existing controls.
- Prioritize remediation efforts and investments based on identified risk.
- Improve detection and response capabilities.
- Maintain continuous improvement in the face of evolving threats.

Agentic AI can accelerate this process by gathering evidence, correlating findings, and linking an identified weakness to possible defensive actions. In this way, the results of offensive assessments do not remain confined to a report, but are turned into practical and measurable improvements.

Conclusion

Offensive and defensive security are not isolated capabilities, nor are they alternatives to one another. Offensive security helps identify and validate weaknesses; defensive security makes it possible to prevent, detect, and respond to them.

Combining both, with the support of agentic AI, enables organizations to build a more proactive, efficient, and adaptable security posture. The goal is not only to react to threats, but to continuously learn from identified weaknesses in order to reduce risk before it materializes.



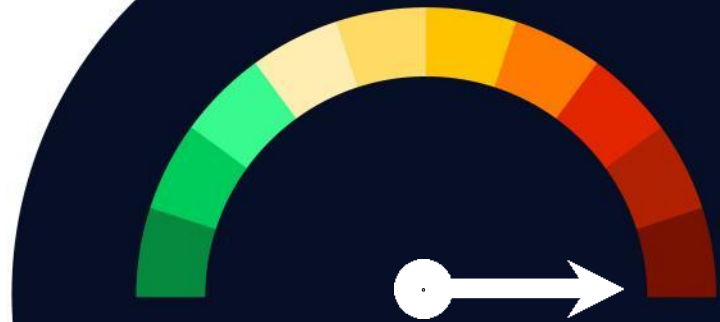
Abel Martin Huaman Condor
Cybersecurity Expert Associate

Vulnerabilities

Veeam ONE – Unauthenticated Remote Code Execution

Date: August 4, 2026

CVE: CVE-2026-64633



CVSS: 10

CRITICAL

Description

A critical vulnerability in Veeam ONE allows an unauthenticated remote attacker to execute arbitrary code on the agent host without user interaction.

Because it requires no credentials and is considered low complexity, it is well suited to automated exploitation, potentially enabling lateral movement, credential theft, or ransomware deployment.

It is especially critical because Veeam ONE monitors backup infrastructure, often the organization's last line of defense, which is why it carries a CVSS score of 10.

Solution

It is recommended to:

- Update to Veeam ONE 13.1.0.7034, released as part of version 13.1, which addresses all the vulnerabilities listed in KB4892.
- There are no alternative mitigations. Until the update is completed, isolate Veeam ONE agent hosts on the network and restrict access to trusted management segments only.

Affected Products

- Veeam ONE 13.0.2.6723 and all earlier builds of version 13.

References

- [veeam.com](https://www.veeam.com)
- cvedetails.com
- [incibe.es](https://www.incibe.es)
- ryzer.es

Vulnerabilities

Multiple Vulnerabilities in VMware Products

Date: July 30, 2026

CVE: CVE-2026-59309 and 4 others



CVSS: 9.8

CRITICAL

Description

VMware has published five vulnerabilities, three of which are rated critical.

The most severe vulnerabilities, identified as CVE-2026-59309, CVE-2026-59310, and CVE-2026-47876, involve an authentication bypass in the VMware Directory service, a directory traversal vulnerability in Syslog, and an out-of-bounds write in the VMXNET3 virtual network adapter.

By exploiting these vulnerabilities, an attacker could bypass authentication and gain unauthorized access to the system or execute arbitrary code on the affected device.

Solution

The vendor recommends updating the affected products to the latest available version.

Affected Products

Some of the affected products include:

- VMware ESX.
- VMware vCenter.
- VMware Workstation.
- VMware Fusion.
- VMware Cloud Foundation.

References

- support.broadcom.com
- incibe.es

Patches

Cisco Security Patch for Critical Catalyst SD-WAN Vulnerabilities

Date: August 5, 2026

CVE: CVE-2026-20303 and 2 others

Critical

Description

Due to insufficient input validation and inadequate access control, vulnerabilities CVE-2026-20303, CVE-2026-20304, and CVE-2026-20310 could allow an attacker to access unintended files and resources, bypass authorization controls, and escalate privileges within the SD-WAN environment.

Catalyst SD-WAN serves as the management plane for corporate connectivity in large organizations. As a result, the potential impact is significant, since a compromise could affect the confidentiality, integrity, and availability of the entire managed network. There are no alternative mitigations to applying the patch.

Affected Products

The affected products are:

- Catalyst SD-WAN Manager, Controller, and Validator in versions prior to the August 2026 hardening releases.

Solution

- Apply the August 2026 hardening releases specified in the Cisco advisory [cisco-sa-hardening-sdwan-faLcR3K](#).
- Restrict access to the management plane to segmented management networks only.

References

- [incibe.es](https://www.incibe.es)
- nvd.nist.gov

Patches

Jenkins Security Update Fixes Critical Vulnerability CVE-2026-70426

Date: August 5, 2026

CVE: CVE-2026-70426 and 21 others

Critical

Description

Jenkins has released a security update addressing a total of 22 vulnerabilities, including one critical vulnerability and eight rated as high severity.

The critical vulnerability, identified as CVE-2026-70426, affects the deserialization mechanism used by the Remoting library and could allow an attacker to send specially crafted serialized objects to execute arbitrary code on the Jenkins controller.

Exploitation of the remaining vulnerabilities could also allow an attacker to bypass security mechanisms, write files to arbitrary locations, or escalate privileges.

Affected Products

Some of the affected products are:

- Jenkins weekly up to version 2.575.
- Jenkins LTS up to version 2.568.1.
- AWS CodeBuild Plugin ≤ 0.59.
- CodeSonar Plugin ≤ 3.6.0.
- External Workspace Manager Plugin ≤ 1.4.1.
- HCL AppScan Plugin ≤ 1.8.3.
- Ivy Report Plugin ≤ 1.2.
- Multijob Plugin ≤ 669.v9d96a_d9c71b_0.
- Parameterized Remote Trigger Plugin ≤ 3.2.2.

Solution

It is recommended to update the affected products to the latest versions released by the vendor.

References

- www.jenkins.io
- www.incibe.es

Events

16th Cloud Security Alliance Spain Meeting

September 17

Next Thursday, September 17, the 16th Cloud Security Alliance Spain Meeting will take place at Kinépolis Ciudad de la Imagen in Madrid. This new edition will focus on the challenges and cloud security threats driven by the rapid evolution of artificial intelligence (AI) and automation.

[Link](#)

27th International Industrial Cybersecurity Congress 2026 (Europe)

September 23-24

Seville will host a new edition of the International Industrial Cybersecurity Congress in Europe on September 23 and 24, bringing together specialists, technology companies, and industry leaders. Organized by the Industrial Cybersecurity Center, the event will focus on the main challenges affecting the protection of critical infrastructure, particularly in a context where industrial environments increasingly combine physical and digital systems.

[Link](#)

CS4CA Europe 2026 — Cyber Security for Critical Assets

September 23-24

The 2026 edition will focus on building intelligence-driven IT/OT resilience. It will address topics such as NIS2, industrial ransomware, remote access, legacy systems, state-sponsored threats, and the protection of sectors including energy, transportation, industry, healthcare, and utilities. The event will feature more than 40 speakers and international IT and OT security leaders.

[Link](#)

International Cyber Expo 2026

September 29-30

It will bring together more than 7,500 professional visitors, representatives from over 90 countries, more than 100 vendors, and more than 100 speakers. The event will include the Global Cyber Summit, hands-on demonstrations, a showcase of security solutions, and participation from CISOs, government officials, regulators, and cybersecurity executives.

[Link](#)

Resources

➤ Trends and Cyber Threats Report

The NTT DATA Cyber Threat Intelligence team's report provides an up-to-date overview of the main threats observed during the first half of 2026, including the evolution of ransomware, malicious actor activity, vulnerability exploitation, dark web trends, the impact of AI on cybercrime, and the main risks facing organizations. It is a key resource for anticipating threats, strengthening defensive capabilities, and guiding cybersecurity decision-making.

[Link](#)

➤ CISA Known Exploited Vulnerabilities Catalog

The Known Exploited Vulnerabilities Catalog, maintained by the U.S. Cybersecurity and Infrastructure Security Agency (CISA), provides an up-to-date list of vulnerabilities for which there is evidence of active exploitation. Unlike other repositories focused solely on technical severity, it enables organizations to prioritize remediation based on real-world risk and attacker activity.

The catalog includes information on the vendor, affected product, required action, and whether the vulnerability may be used in ransomware campaigns. It can also be consulted online or downloaded in CSV and JSON formats for integration into vulnerability management processes and tools.

[Link](#)

➤ M-Trends 2026: Threats Observed from the Front Lines

The M-Trends 2026 report, produced by Mandiant — Google Cloud — examines the evolution of the global cyber threat landscape based on more than 500,000 hours of investigations and incident response work carried out during 2025. The report covers topics such as attackers' use of artificial intelligence, the increasing professionalization of the ransomware ecosystem, initial access techniques, persistence on edge devices, and the need to evolve from reactive defense toward active resilience models.

It is an especially valuable resource for SOC teams, incident response professionals, threat intelligence teams, and cybersecurity strategy leaders.

[Link](#)



Subscribe to RADAR

up.nttdata.com/suscribetearadar

**Powered by the
cybersecurity
NTT DATA team**

es.nttdata.com