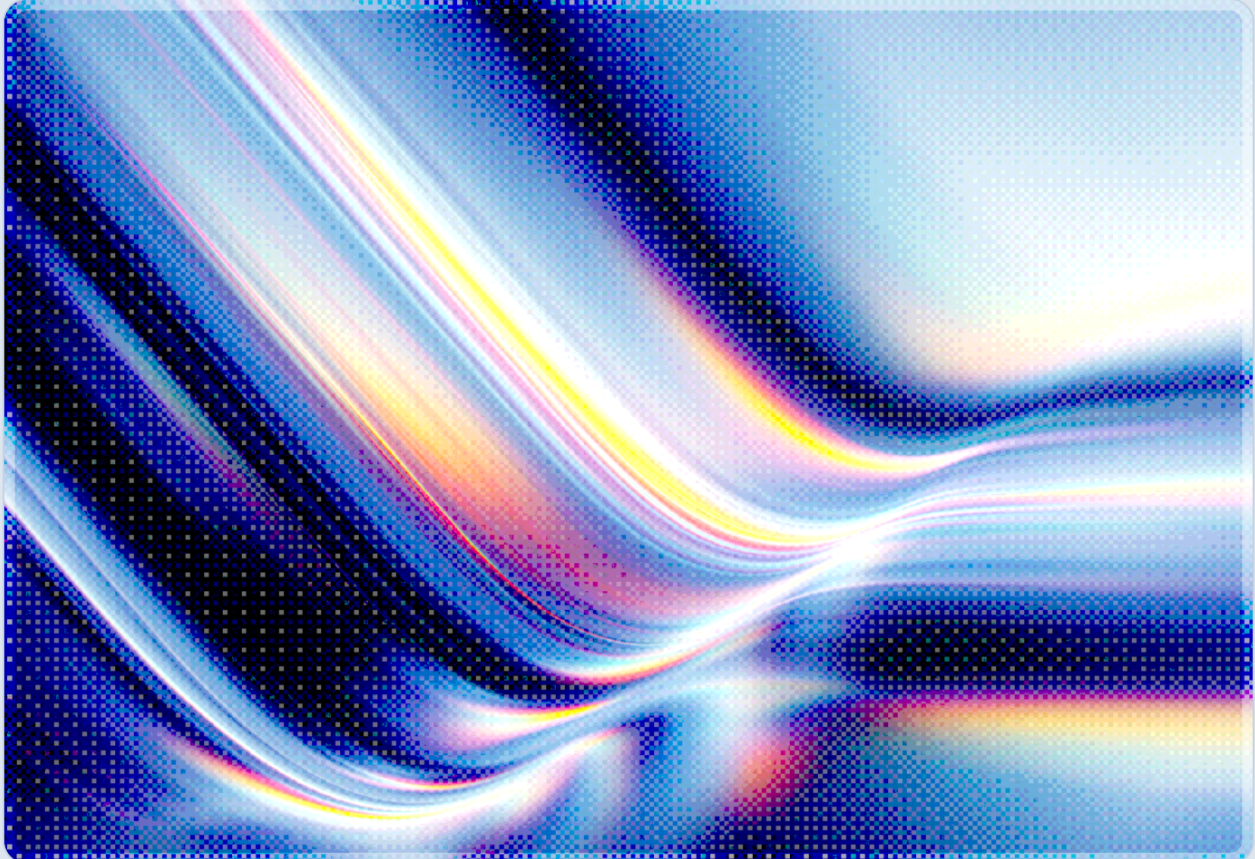


DeepSeek-Driven Autonomous Cyberattacks: A Solo Operator's Campaign

Hermes Agent, FOFA, and the Machine-Speed Exploitation of Langflow, n8n, and Citrix NetScaler

2026-08-04

 AI-assisted Rapid Research



CSAI

cloud
security
alliance®

© 2026 Cloud Security Alliance. Some rights reserved.

You may download, store, display, view, print, redistribute, and link to this document in its original, unmodified form, provided that attribution to the Cloud Security Alliance is maintained and all trademark and copyright notices remain intact.

This document may not be modified or altered. You may quote portions of the document as permitted by the Fair Use provisions of the United States Copyright Act, provided that attribution is given to the Cloud Security Alliance.

This document may be shared on professional and social media platforms in its original form with attribution.

This document was generated with AI assistance and has not undergone official CSA review and approval processes.

Key Takeaways

A China-based individual operating under the aliases "knaithe" and "KnYuan" configured DeepSeek as the reasoning engine inside Hermes Agent, an open-source terminal-access agent framework from Nous Research, and directed it to conduct offensive operations against internet-facing servers with minimal human involvement, according to a Palo Alto Networks Unit 42 disclosure [1]. The operator issued instructions over Telegram and enabled Hermes Agent's "Yolo" mode, which allows the agent to execute commands, including destructive or risky ones, without requesting approval, and the agent then independently enumerated vulnerable infrastructure, researched candidate CVEs, downloaded proof-of-concept exploit code, and attempted exploitation across seven distinct attack tracks spanning eight CVEs [1][2].

The campaign, reconstructed from a recovered May 2026 session, targeted more than 460 systems, including 84 exposed Langflow instances and a pool of 647,017 internet-facing n8n workflow-automation deployments worldwide, of which 25,209 were located in China [1]. Unit 42 reported that the autonomous system "executed hundreds of hours of manual targeting analysis in mere minutes," compressing reconnaissance and triage work that would otherwise require sustained human effort [1][3]. Despite that scale, the operation's confirmed successful compromises were limited to three organizations, illustrating that autonomous reconnaissance and triage scaled far more readily than autonomous exploitation [1].

Confirmed impact was concentrated in two of the eight tracked vulnerabilities: a Citrix NetScaler flaw (CVE-2026-3055, CVSS 9.8) yielded data exfiltration from three targets, after which the actor searched the stolen data for NSC_AAAC session-authentication cookies to enable session hijacking, and a Marimo notebook vulnerability (CVE-2026-39987, CVSS 9.8) yielded confirmed command execution on eleven targets [1]. Autonomous attempts against Langflow (CVE-2026-33017) and an n8n vulnerability chain (CVE-2026-21858 and CVE-2025-68613) failed outright due to missing deployment prerequisites, and other tracks produced only reverse-shell attempts or non-functional proof-of-concept clones [1].

The operation was discovered only because Hermes Agent's default behavior undermined the operator's own security discipline: the framework started an HTTP file server from the actor's home directory, exposing API keys, exploit scripts, target lists, shell history, and full session logs to researchers, even though the actor had otherwise practiced deliberate operational security by emptying exploit directories after use and disabling conversation logging in a secondary tool [1]. Unit 42's disclosure identifies no state-sponsorship indicators, and the resourcing and scope evident in the account are consistent with a

single individual rather than a state-backed team assembling a largely autonomous offensive workflow from a commodity reasoning model and a freely available open-source agent framework – a combination CSA's own research has identified as a defining feature of the current threat landscape [4][5][6].

Background

Hermes Agent is a legitimate, actively maintained open-source project published by Nous Research that provides terminal access, a reusable skills system, and integrations with more than twenty large language model providers, allowing an operator to route the same agentic workflow through DeepSeek, Anthropic, OpenAI, Google, xAI, or a locally hosted model interchangeably [7]. In this campaign, the actor configured Hermes Agent to use DeepSeek as its primary reasoning backend and connected the framework to Telegram for command-and-control, giving the operator a lightweight mobile interface for issuing high-level instructions while the agent handled tactical execution [1][2]. The actor also experimented with several alternative backends during the same period, including limited connectivity testing of Claude Code, minimal use of OpenAI's Codex on exploit-development directories, and evaluation of Chinese-market models Qwen, GLM, Kimi, and MiniMax, a pattern consistent with an operator treating the choice of reasoning model as a fungible, swappable component of the offensive pipeline rather than a fixed dependency [1].

The actor's public presence, including a GitHub profile and an older personal blog describing the author as a Zhuhai-based binary security researcher, is also linked to 1DayNews, an automated pipeline that aggregates newly disclosed remote-code-execution vulnerabilities from seventeen sources [1]. That existing infrastructure fed directly into the autonomous campaign: the agent was equipped with custom skills named "fofa-cyberspace-search" and "web-terminal-exploitation," and the operator integrated a Model Context Protocol server called FofaMap-Platinum-Full-Expert that translated natural-language requests into structured queries against FOFA, a Chinese internet-asset search engine comparable to Shodan or Censys [1]. This combination let the agent move from a general instruction to a concrete list of candidate targets without the operator manually crafting search-engine syntax, a step that has traditionally required specialized reconnaissance skill.

The attack cycle documented in the recovered session followed a repeatable pattern. The agent first queried FOFA to enumerate internet-facing instances of a candidate product, then, when an initial exploitation attempt failed, independently surveyed roughly ten additional product families, searched GitHub for trending proof-of-concept code, and prioritized remaining candidates by attack surface and CVSS severity before downloading exploit code and resuming testing [1]. For the n8n track specifically, the agent narrowed a population of 25,209 Chinese-based instances down to a sample of roughly one hundred, actively probed about forty unique IP addresses, and identified three that were running a

vulnerable version, all without operator intervention at each intermediate step [1]. This is the behavior Unit 42 characterized as compressing what would otherwise be hundreds of hours of manual targeting analysis into minutes, and it is the clearest evidence in the disclosure that the autonomy extended meaningfully beyond simple script execution into iterative research and prioritization [1][3].

Security Analysis

The most consequential finding in this disclosure is the sharp divergence between what the autonomous system accomplished on its own and what still required manual follow-through. Target discovery, CVE research, proof-of-concept acquisition, and initial exploitation attempts were handled by DeepSeek operating inside Hermes Agent with limited operator input. Confirmed, weaponizable outcomes – sustained data exfiltration from the Citrix NetScaler targets and the multi-day, persistent targeting of a Malaysian government entity – depended on custom Python scanners and direct human-directed exploitation layered on top of the agent's reconnaissance [1]. Read against the three confirmed compromises out of more than 460 attempted targets, the case supports a measured reading of current autonomous offensive capability: it substantially lowers the cost of reconnaissance, triage, and opportunistic probing at scale, but exploitation against patched, well-configured, or otherwise unfavorable targets still routinely fails without a human closing the gap. That same asymmetry between machine-speed reconnaissance and comparatively brittle autonomous exploitation is the central empirical finding of CSA's broader survey of documented autonomous adversary cases, including the state-sponsored GTG-1002 campaign and the CyberStrikeAI operation against internet-facing firewalls [5].

The disclosure also surfaces what can be described as an operational security paradox specific to agentic tooling. The actor demonstrated conventional tradecraft discipline by clearing exploit directories after use and disabling logging in a secondary tool, yet Hermes Agent's own default behavior of exposing a directory over HTTP for convenience undid that discipline in a single step, handing Unit 42 the operator's API keys, target lists, and session history [1]. This is a variant of a pattern CSA's research on LLM-driven post-exploitation has already documented: agent frameworks optimized for operator convenience and rapid iteration frequently carry default behaviors – verbose logging, permissive file-serving, broad tool access – that are misaligned with the operational security an offensive user would otherwise want, and that misalignment is currently a meaningful, if temporary, source of defensive advantage [4]. Attackers who adopt agent frameworks inherit not only their capabilities but also their defaults, and the defaults are set by developers building tools for legitimate automation and coding use cases, not for operational security under adversarial scrutiny.

Finally, the campaign is notable for what it implies about the accessibility of autonomous offensive capability rather than for any single technical novelty. Every component the actor assembled – DeepSeek's API, the open-source Hermes Agent framework, the FOFA search engine, and a Telegram bot for command relay – is either free, low-cost, or available to any individual with modest technical skill. CSA's research into automated exploit generation has already documented that the cost of moving from a disclosed vulnerability to working exploit code has fallen substantially as reasoning models have matured, and this campaign is a concrete instance of that trend applied end to end, by one person, against a global pool of exposed workflow-automation and notebook platforms [6]. The gap between this case and a nation-state operation like GTG-1002 is not the presence or absence of autonomy; it is scale, resourcing, and the sophistication of the human decision-making layered on top of a broadly similar autonomous substrate [5].

Recommendations

Immediate Actions

Organizations running Langflow, n8n, Marimo notebooks, Citrix NetScaler, Apache Tomcat, or systems exposing the Windows IKE VPN extension or PAN-OS captive portal should confirm patch status against the specific CVEs implicated in this campaign – CVE-2026-33017, CVE-2026-21858, CVE-2025-68613, CVE-2026-3055, CVE-2026-39987, CVE-2026-34486, CVE-2026-33824, and CVE-2026-0300 – and treat any unpatched, internet-facing instance as an active exposure rather than a theoretical one [1]. Organizations operating Citrix NetScaler appliances that may have been exposed to CVE-2026-3055 should rotate NSC_AAAC session-authentication cookies and review authenticated sessions for anomalous activity, given the actor's documented practice of searching exfiltrated data specifically for that cookie to enable session hijacking [1]. Any organization running Marimo notebooks with confirmed or suspected exposure to CVE-2026-39987 should assume command execution may have occurred and initiate incident-response triage on the affected hosts.

Short-Term Mitigations

Security teams should treat internet-facing workflow-automation platforms and notebook environments – Langflow and n8n in particular, given their combined exposed footprint of hundreds of thousands of instances – as a distinct, high-priority asset class for external attack-surface management, since these platforms often hold embedded API keys and credentials for the services they orchestrate and are attractive both to autonomous scanning and to opportunistic human attackers [1]. Network detection should be tuned to flag the behavioral signature of agent-driven reconnaissance: rapid, sequential

probing of many distinct IP addresses at a cadence inconsistent with manual operation, and outbound connections consistent with search-engine APIs such as FOFA, Shodan, or Censys originating from unexpected internal hosts. Egress monitoring should also flag outbound Telegram Bot API traffic from server-side infrastructure, which is not a typical pattern for legitimate enterprise workloads and was the actor's chosen command-and-control channel in this case [1].

Strategic Considerations

In CSA's assessment, this campaign is evidence that autonomous offensive tooling assembled from commodity reasoning models and freely available agent frameworks now sits within reach of a single individual of modest resources, not only well-resourced state actors, and organizations should adjust their threat models accordingly to anticipate a higher volume of similarly structured, broadly scoped, opportunistic campaigns [1][5]. The operational security failures that led to this campaign's discovery reflect defaults in current agent tooling rather than a durable defensive advantage, and defenders should not assume that future autonomous offensive operators will make the same mistake; detection strategies should be built around the behavioral signatures of the attack pattern itself rather than around the assumption that the tooling will continue to expose itself. Finally, the wide gap between attempted targets and confirmed compromises in this case is a reminder that patch currency and secure default configuration remain the most effective controls against both human and autonomous adversaries, since the vulnerabilities that produced confirmed impact were, in each case, ones where the target environment had not been updated or hardened against a publicly known flaw.

CSA Resource Alignment

CSA Labs' research note on the Marimo remote-code-execution vulnerability is the most directly applicable prior CSA work: it documents CVE-2026-39987 – the same vulnerability responsible for the eleven confirmed command-execution outcomes in this campaign – and catalogs a separate, previously observed in-the-wild intrusion in which an LLM agent used Marimo access to pivot through a target environment [4]. That research identifies behavioral fingerprints distinguishing LLM-agent-driven post-exploitation from human-operated intrusions, including machine-optimized command syntax and output-dependent value chaining, and those same fingerprints are a useful starting point for detection engineering aimed at the DeepSeek and Hermes Agent workflow described here.

CSA's whitepaper on automated exploit generation provides the analytical framework for understanding why this campaign was economically viable for a solo operator [6]. That research tracked the decline in cost and skill required to move from a disclosed vulnerability to working exploit code as reasoning models

matured, and the DeepSeek-driven campaign's autonomous CVE triage, GitHub proof-of-concept retrieval, and rapid pivoting across ten product families is a direct, real-world instance of the trajectory that paper anticipated.

CSA's whitepaper on autonomous agentic AI adversaries situates this campaign within a broader, growing set of documented cases in which frontier or near-frontier AI models have been deployed as autonomous operational actors rather than passive tools, alongside the state-sponsored GTG-1002 campaign and the CyberStrikeAI operation against internet-facing firewalls [5]. That paper's mapping of autonomous-adversary behavior to CSA's MAESTRO agentic AI threat-modeling framework and the AI Controls Matrix (AICM) gives defenders a structured way to fold this new case into existing agentic AI risk assessments rather than treating it as an isolated incident.

References

- [1] Palo Alto Networks Unit 42. "[Chinese-Speaking Threat Actor Harnesses AI Models for Autonomous Cyberattacks](#)." Unit 42, July 30, 2026.
- [2] The Hacker News. "[Chinese Hacker Commands DeepSeek via Telegram to Launch Autonomous Attacks](#)." The Hacker News, July 31, 2026.
- [3] BleepingComputer. "[Hacker Uses DeepSeek AI to Autonomously Attack Vulnerable Servers](#)." BleepingComputer, July 31, 2026.
- [4] Cloud Security Alliance Labs. "[Marimo RCE: LLM Agents as Post-Exploitation Tools](#)." CSA Labs, June 6, 2026.
- [5] Cloud Security Alliance Labs. "[Autonomous Agentic AI Adversaries](#)." CSA Labs, June 24, 2026.
- [6] Cloud Security Alliance Labs. "[Automated Exploit Generation: LLMs Cross the Threshold](#)." CSA Labs, April 2, 2026.
- [7] Nous Research. "[Hermes Agent Documentation](#)." Nous Research, 2026.