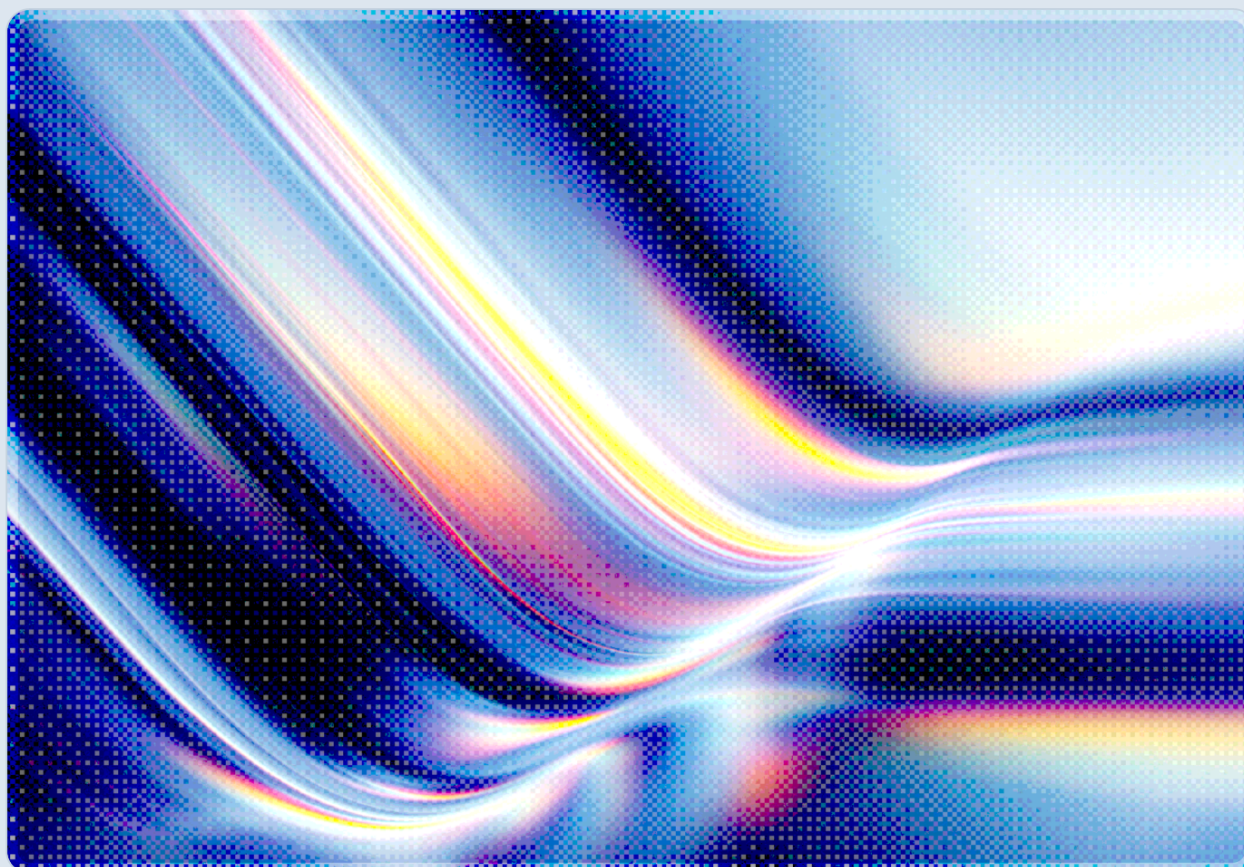


AI Infrastructure Is Now a Standing Attacker Target Class

2026-08-31

 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

Ninety days of honeypot telemetry published by Wiz Research show that self-hosted AI tooling – model gateways, orchestration frameworks, vector databases, and inference servers – is now attacked with the same persistence and tooling maturity that commodity web infrastructure has faced for two decades [1]. The exposure is not a hypothetical: Wiz's own 2026 cloud survey found that 90 percent of cloud environments run self-hosted AI software and that Model Context Protocol (MCP) servers now appear in roughly 80 percent of observed environments, with 5 percent of those exposing an MCP server directly to the internet [2]. Attackers have moved beyond opportunistic credential theft toward AI-native post-exploitation techniques, including querying live process memory for API keys and using blind prompt injection with out-of-band DNS callbacks to confirm code execution inside orchestration frameworks [1]. Two LiteLLM vulnerabilities – an MCP gateway authentication bypass (CVE-2026-59822) and an MCP test-endpoint command injection (CVE-2026-42271) that CISA added to its Known Exploited Vulnerabilities catalog – were both observed being exploited in the wild inside the honeypot window, and the latter can be chained with a Starlette authentication bypass for unauthenticated remote code execution rated CVSS 10.0 [1][3][4]. Wiz separately reports that attacker activity targeting AI infrastructure roughly doubled between the second half of 2025 and the first half of 2026, and that LiteLLM alone – present in more than a third of monitored cloud environments – suffered four distinct security incidents in six months [5]. Organizations should treat every self-hosted AI component as internet-facing infrastructure by default: inventory it, authenticate it, segment it from cloud credentials, and monitor its process behavior, rather than continuing to manage it as internal developer tooling.

Background

Security teams have widely, though not universally, treated self-hosted AI tooling – LangChain pipelines, Ollama model servers, vector databases such as ChromaDB, low-code agent builders such as Flowise and Langflow – as internal developer convenience software: fast to stand up, rarely inventoried, and rarely subjected to the change-management or network-exposure review applied to production services. CSA's own prior research on exposed Ollama and MCP infrastructure is consistent with this pattern: rapid-research threat notes have already documented a botnet that chains more than 20 exploitation techniques specifically to harvest cloud keys from exposed AI infrastructure [6], and a separate incident in which an exposed, unauthenticated Ollama server was repurposed as the reasoning engine for a nine-

stage autonomous exploitation framework [7]. That posture toward AI tooling assumed attackers had not yet developed the tradecraft or motivation to target it, an assumption the evidence accumulated over 2026 undercuts. Wiz's 90-day honeypot study extends this picture from isolated incidents to sustained, structural targeting: the firm operated honeypots emulating LiteLLM, Flowise, LangChain, Langflow, ChromaDB, Ollama, and related services, and observed attackers using purpose-built tradecraft against each one rather than generic internet-scanning noise [1].

One reason this shift matters is structural: AI infrastructure often occupies a privileged position in the modern cloud stack. By design, a LiteLLM gateway or MCP server generally needs API keys to one or more upstream model providers and broad IAM roles or service-account tokens for the workloads it orchestrates, which is what makes exposure consequential rather than incidental; it frequently also carries network reachability into the rest of a cluster, reflecting deployment patterns that did not anticipate the service becoming an entry point. Wiz's broader 2026 cloud survey found that 81 percent of cloud environments run managed AI services and 90 percent run self-hosted AI software, that 68 percent of organizations ingest models through third-party software (with 18 percent relying exclusively on such "transitive" dependencies), and that MCP servers now appear in roughly 80 percent of observed environments [2]. That scale of adoption, combined with security practices that Wiz characterizes as not having "caught up to how widely" these tools are deployed, is consistent with the conditions under which a target class forms: a large, unauthenticated-by-default, credential-rich population of services that is straightforward to fingerprint and exploit at internet scale [5].

Security Analysis

The honeypot telemetry identifies three attack patterns that, taken together, describe how attackers now operationalize access to exposed AI infrastructure rather than merely stumbling onto it.

MCP server exploitation using named CVEs. LiteLLM's MCP Gateway was hit through two distinct vulnerability classes inside the observation window. CVE-2026-59822 is an authentication-bypass flaw in the MCP Streamable HTTP endpoint: when LiteLLM key validation fails, the server's OAuth2 passthrough fallback path returns an empty, unrestricted authorization object, meaning any Bearer token – even a single character – grants full access to configured MCP tools and the backend services they connect to [3]. CVE-2026-42271 is a command-injection flaw in LiteLLM's MCP test endpoints (`/mcp-rest/test/connection` and `/mcp-rest/test/tools/list`), which accept attacker-supplied `command`, `args`, and `env` fields for stdio-transport MCP servers and pass them to the host for execution, giving an authenticated (even low-privilege) caller command execution with proxy-process privileges [4]. CISA added CVE-2026-42271 to its Known Exploited Vulnerabilities catalog on June 8, 2026, roughly five weeks after a patch became available [1][4][8], and Horizon3.ai

subsequently demonstrated that chaining it with a Starlette Host-header authentication bypass (CVE-2026-48710, "BadHost") removes the authentication requirement entirely, producing a combined CVSS 10.0 unauthenticated remote-code-execution path [1][4]. Wiz's honeypots recorded both vulnerabilities being used to deliver Python-based cryptominers that auto-delete their own staging directories after execution, an evasion behavior consistent with an effort to frustrate post-incident forensics [1].

Blind prompt injection as a delivery mechanism. A second pattern targeted LangChain, Flowise, OpenWebUI, and Node-RED deployments not through a software vulnerability but through the applications' own designed behavior: injecting prompts engineered to trigger shell command execution within the orchestration layer. Because these frameworks frequently lack verbose execution logging, attackers confirmed success out-of-band, via DNS callbacks to attacker-controlled domains, and then staged payloads as base64-encoded strings fetched from Pastebin, a technique consistent with an effort to avoid leaving plaintext commands in application logs, ultimately deploying XMRig cryptomining payloads [1]. This is a materially different exploitation model than a CVE: it works against fully patched software because it abuses legitimate prompt-to-tool-execution functionality, and is unlikely to be resolved by a vendor patch alone.

AI-native post-exploitation. The most notable shift in the telemetry is what attackers did after gaining a foothold. Rather than the traditional approach of grepping the filesystem for credential files, attackers queried running Python processes directly to extract LiteLLM's in-memory `master_key` value – a technique that presumes attacker familiarity with the internal architecture of the specific AI framework being targeted, not generic post-exploitation tradecraft [1]. Attackers also systematically enumerated framework-specific configuration paths and fingerprinted which backend models were reachable using default credentials, behavior consistent with reconnaissance aimed at building a reusable playbook against each AI framework rather than a one-off intrusion [1]. This mirrors the pattern CSA documented in the LLMjacking evolution research: attackers are treating exposed AI infrastructure not merely as a resource to steal, but as a substrate – reasoning engine, compute, or credential store – for further offensive activity [7].

The scale context from Wiz's 2026 cloud survey helps explain why this tradecraft investment pays off. LiteLLM alone is present in more than a third of monitored cloud environments and experienced four separate security incidents in six months, spanning a supply-chain compromise, an in-the-wild SQL injection, a privilege-escalation chain, and the authentication-bypass flaws described above; critical unauthenticated vulnerabilities were also identified in Dify, Langflow, n8n, and Ollama over the same period [5]. Wiz reports that overall attacker activity targeting AI infrastructure roughly doubled from the second half of 2025 to the first half of 2026 – a trajectory consistent with a target class still in its early growth phase rather than one that has plateaued [5]. CSA's own threat research corroborates the exposure side of this equation: an earlier CSA note cited approximately 175,000 publicly exposed Ollama instances across more than 130 countries, most bound to all network interfaces without authentication

by default [7], and CSA's NadMesh analysis documented a single botnet chaining more than 20 remote-code-execution vectors – including several against AI-specific tools such as ComfyUI and Ollama – that reported harvesting several thousand cloud credential sets in a single week of operation [6]. None of these figures should be treated as independently audited; Wiz's honeypot post does not disclose raw attack counts, and CSA's own prior notes caution that attacker-reported operational figures are directional rather than verified. But the consistent direction across four independently sourced datasets – honeypot telemetry, a cloud-wide adoption survey, a botnet campaign analysis, and a live exploitation incident – is what elevates this from an isolated finding to a structural conclusion: AI infrastructure has become a standing target class, attacked continuously and with increasingly specialized tradecraft, rather than a source of occasional, generic opportunistic hits.

Recommendations

Immediate Actions

Security and platform teams should confirm, within days rather than weeks, whether any LiteLLM, Ollama, Flowise, LangChain, Langflow, ChromaDB, or similar self-hosted AI service is reachable from the public internet, and if so, place it behind authentication and network access controls or take it offline until it can be properly secured. Any LiteLLM deployment should be upgraded to version 1.84.0 or later to remediate CVE-2026-59822, and to version 1.83.7 or later (with Starlette upgraded to 1.0.1 or later) to remediate CVE-2026-42271 and its BadHost chain; where immediate upgrade is not possible, MCP routes should be disabled or blocked at the reverse proxy [3][4]. All API keys and tokens reachable from an exposed AI service – including LLM provider keys, cloud credentials, and Kubernetes service-account tokens – should be treated as potentially compromised and rotated, with usage logs reviewed for anomalous activity predating rotation.

Short-Term Mitigations

Organizations should complete an honest inventory of self-hosted AI tooling deployed outside formal change management, since much of this infrastructure was originally stood up by data science or engineering teams without security review. Every AI gateway, model server, and orchestration framework identified should be bound to loopback or an internal network segment rather than a public interface, fronted by mandatory authentication, and instrumented with full request and response logging fed into a SIEM with correlation rules tuned for the AI-native behaviors described above – in-memory key extraction, unusual MCP tool invocation patterns, and outbound DNS requests to unfamiliar domains

immediately following an inbound prompt. Patch cadence for open-source AI development tools should be aligned with, not deprioritized relative to, core infrastructure patching, given that CVE-2026-42271 reached CISA's KEV catalog within five weeks of patch availability.

Strategic Considerations

Over the medium term, organizations should extend existing identity-and-access-management and zero-trust architecture programs to explicitly cover AI infrastructure rather than treating it as an exception. This includes issuing short-lived, scoped credentials to AI services instead of long-lived static keys, applying default-deny network posture to newly deployed AI tooling, and building architecture-review checkpoints that require network binding, authentication, and logging decisions to be made explicitly before any new AI framework goes into production. Given the doubling in attacker activity that Wiz documented over a single six-month period, security leadership should plan for AI infrastructure exposure reviews to become a recurring, not one-time, control – comparable in cadence to how organizations already treat internet-facing web application scanning.

CSA Resource Alignment

This research connects directly to three pieces of CSA's own recent threat intelligence work on AI infrastructure exposure. CSA's rapid-research note on the [LiteLLM AI Gateway command-injection campaign](#) analyzed CVE-2026-42271 in detail at the time CISA added it to the KEV catalog, and its recommendations on upgrade timelines, credential rotation, and MCP endpoint monitoring apply directly to the exploitation of that same flaw documented in the Wiz honeypot telemetry. CSA's threat note on [NadMesh, a botnet built to hunt exposed AI infrastructure for cloud keys](#) [6] documents the same underlying dynamic from the attacker-tooling side – a single exposed AI instance functioning as an entry point into broader cloud and cluster compromise – and its recommendation to apply a default-deny posture to AI tooling exposure is reinforced by the honeypot findings here. CSA's analysis of [LLMjacking's evolution into offensive AI infrastructure](#) [7] documented attackers using an exposed, unauthenticated Ollama server as a reasoning engine for autonomous exploitation, establishing the AI-native post-exploitation pattern that the Wiz telemetry now shows generalizing across the broader self-hosted AI ecosystem. Finally, the identity, credential-rotation, and network-segmentation recommendations throughout this note map to the Identity and Access Management, Application and Infrastructure Security, and Threat and Vulnerability Management domains of CSA's [AI Controls Matrix \(AICM\) v1.1](#) [9], which organizations should use as the control baseline when formalizing governance over self-hosted AI infrastructure.

References

- [1] Wiz Research. "[Attacks on AI Infrastructure: 90-Day Honey,pot Telemetry.](#)" Wiz Blog, 2026.
- [2] Wiz Research. "[State of AI in the Cloud 2026.](#)" Wiz, 2026.
- [3] GitHub / BerriAI. "[MCP Authentication Bypass via OAuth2 Passthrough Fallback \(CVE-2026-5982 2\).](#)" GitHub Security Advisories, 2026.
- [4] Horizon3.ai. "[CVE-2026-42271: LiteLLM Unauthenticated RCE Chained with CVE-2026-48710.](#)" Horizon3.ai Attack Research, 2026.
- [5] Wiz Research. "[Cloud Threat Highlights: H1 2026.](#)" Wiz Blog, 2026.
- [6] Cloud Security Alliance. "[NadMesh: A Botnet Built to Hunt Exposed AI Infrastructure for Cloud Keys.](#)" Cloud Security Alliance, July 2026.
- [7] Cloud Security Alliance. "[LLMjacking Evolved: Stolen AI Compute as Offensive Infrastructure.](#)" Cloud Security Alliance, June 2026.
- [8] The Hacker News. "[LiteLLM Flaw CVE-2026-42271 Exploited in the Wild, Chains to Unauthenticated RCE.](#)" The Hacker News, June 2026.
- [9] Cloud Security Alliance. "[AI Controls Matrix \(AICM\) v1.1.](#)" Cloud Security Alliance, 2026.