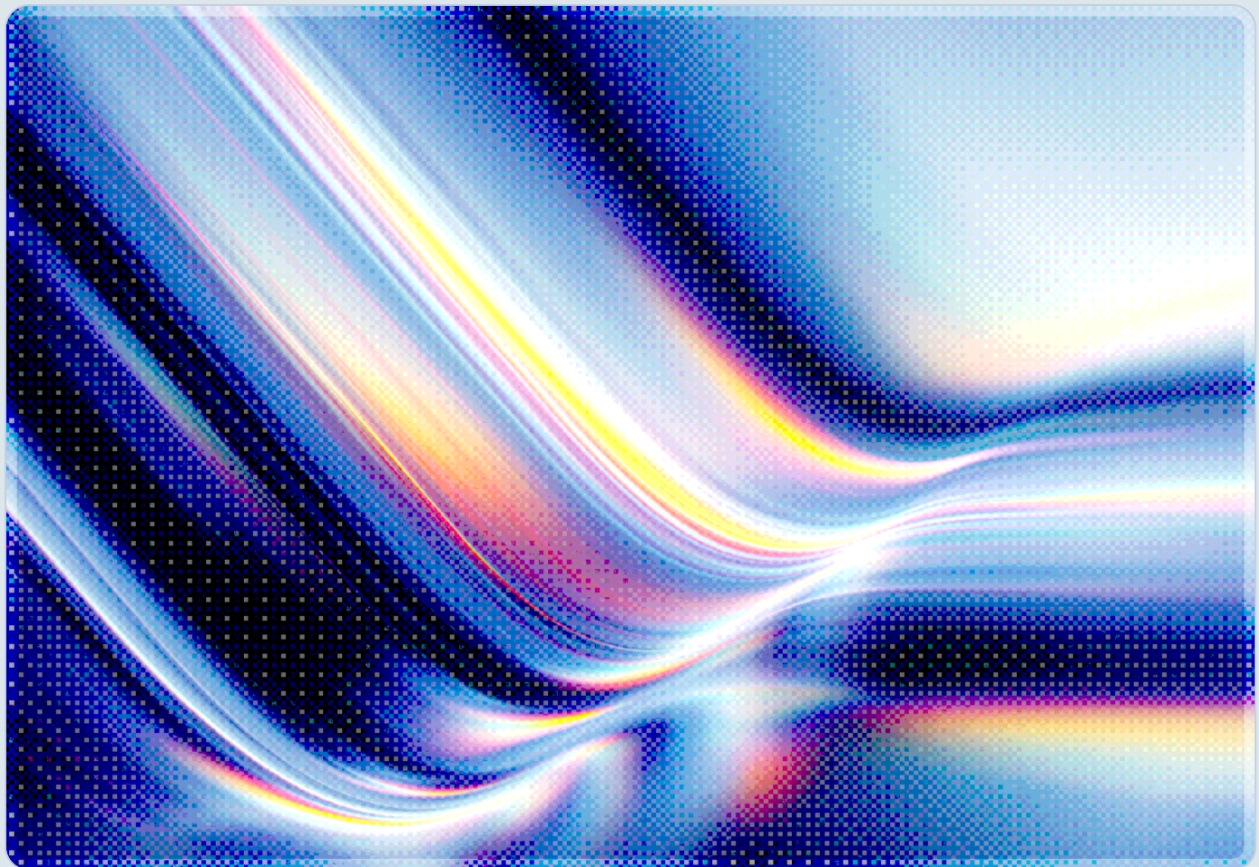


Langflow Auto-Login Bypass Chains to Unauthenticated RCE

2026-08-18

 AI-assisted Rapid Research



CSAI

cloud
CSA 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

CVE-2026-9198 is a critical (CVSS 9.8) vulnerability chain in the open-source Langflow AI orchestration platform that lets a fully unauthenticated attacker obtain a superuser session token and then use it to execute arbitrary Python code on the host, with no login and no user interaction required [1][2]. The chain combines two component flaws, CVE-2026-9103 and CVE-2026-8481, in the `/api/v1/auto_login` and `/api/v1/validate/code` endpoints, respectively [3][4]. The vulnerability affects Langflow OSS versions 1.0.0 through 1.10.0 and was fixed in version 1.10.1, released the same day IBM disclosed it, July 17, 2026 [2][3]. A public proof-of-concept began circulating in late July, and CISA added CVE-2026-9198 to its Known Exploited Vulnerabilities (KEV) catalog on August 4, 2026, after confirming active exploitation [1][5]. Telemetry gathered since early July shows sustained scanning and exploitation from hundreds of distinct source IP addresses across dozens of countries, and prior Langflow incidents in 2026 show that attackers who gain this level of access typically move quickly to harvest LLM provider keys, cloud credentials, and database secrets embedded in workflow configurations [6][7]. Organizations running any internet-reachable Langflow instance at or below version 1.10.0 should treat this as an emergency patching event and assume compromise if the instance has been exposed since early July.

Background

Langflow is a widely adopted open-source, low-code platform for building and orchestrating large language model applications and AI agent workflows; it has attracted more than 150,000 GitHub stars and is distributed commercially in part through IBM's involvement in the project [2][6]. Its visual, drag-and-drop interface for chaining LLM calls, retrieval-augmented generation components, and custom Python logic has made it popular with teams standing up AI proofs-of-concept and production pipelines alike, but that same flexibility means a compromised Langflow instance often sits at the center of an organization's AI credential surface, holding API keys for model providers, database connection strings, and cloud service credentials used by the flows it runs [6][7].

CVE-2026-9198 is the fourth distinct unauthenticated or authorization-bypass vulnerability in Langflow to draw CISA KEV attention within roughly a five-month span in 2026, following a build-endpoint RCE flaw (CVE-2026-33017, KEV-listed in March), a file-upload path traversal flaw (CVE-2026-5027, disclosed in June), and a cross-tenant authorization bypass (CVE-2026-55255, KEV-listed in July) [6]

[8]. Each of these incidents traces back to a similar root pattern: default configurations or endpoint logic that assumes network-level trust rather than enforcing authentication and authorization at the API layer, a pattern common to AI development tooling that was built for rapid iteration rather than production-grade hardening [6][7]. CVE-2026-9198 fits squarely into this pattern and, notably, involves an entirely separate code path from the three prior disclosures, indicating the underlying authentication and code-execution model in Langflow has required repeated, piecemeal remediation rather than a single structural fix [3][4].

The specific defect at the center of CVE-2026-9198 involves the `/api/v1/auto_login` endpoint, which is intended to streamline local development by automatically issuing a session token, but which minted a SUPERUSER-scoped bearer token to any network caller regardless of whether authentication credentials were supplied [2][3]. This behavior is tracked separately as CVE-2026-9103 [4]. IBM disclosed both the auto-login flaw and the second half of the chain on July 17, 2026, and shipped Langflow 1.10.1 the same day [2][3]. The current stable release as of this writing is 1.11.2 [11].

Security Analysis

The second half of the exploit chain targets the `/api/v1/validate/code` endpoint, a feature that lets Langflow evaluate user-submitted Python snippets, for example to validate custom component logic before it is added to a workflow. That endpoint passes attacker-supplied input directly into Python's `exec()` function rather than parsing or sandboxing it, a defect tracked as CVE-2026-8481 [3][4]. On its own, this endpoint should require an authenticated session; the auto-login flaw removes that barrier entirely. An attacker needs only two HTTP requests: the first, an unauthenticated call to `/api/v1/auto_login`, returns a valid superuser JSON Web Token; the second, a POST to `/api/v1/validate/code` bearing that token and a malicious Python payload, executes with the full privileges of the Langflow service process [2][3]. Because the flow requires no credentials, no prior reconnaissance beyond confirming the target is running Langflow, and no user interaction, CISA and multiple vendor advisories rate it CVSS 9.8, near the ceiling of the scoring system [1][2].

Exploitation telemetry corroborates that this is not a theoretical risk. Independent monitoring recorded exploitation attempts beginning around July 6, 2026, weeks before IBM's formal disclosure and patch, which suggests either independent discovery by opportunistic scanners or early leakage of technical details ahead of coordinated disclosure [5]. By mid-August, cumulative attempt counts had grown into the hundreds, originating from several hundred unique attacker IP addresses spread across dozens of countries, a pattern consistent with both automated internet-wide scanning and more targeted follow-on activity [5]. Public proof-of-concept exploit code appeared on code-sharing platforms in late July,

which typically accelerates both the volume and geographic diversity of opportunistic attacks, and researchers monitoring the campaign describe exploit-development artifacts and staging infrastructure consistent with active weaponization rather than isolated curiosity scanning [1][9].

The practical consequences of a successful chain closely mirror what has already been documented in Langflow's two prior 2026 RCE incidents. Attackers who achieve code execution on a Langflow host gain direct access to whatever secrets the platform has been configured to hold on behalf of its workflows, commonly including OpenAI, Anthropic, or other LLM provider API keys, cloud service credentials, database connection strings, and any documents or embeddings ingested into connected retrieval-augmented generation stores [6][7]. Because Langflow instances are frequently deployed by individual teams or managed service providers on behalf of multiple downstream customers, compromise of a single instance can expose credentials and data belonging to several unrelated organizations, and prior campaigns against this platform have deployed cryptocurrency-mining payloads once access is established [1][7]. The technical simplicity of the CVE-2026-9198 chain, combined with the value of what a compromised instance typically exposes, makes it an attractive and low-cost target for both financially motivated actors and more patient adversaries interested in credential harvesting and supply chain positioning.

CISA's decision to add CVE-2026-9198 to the Known Exploited Vulnerabilities catalog on August 4, 2026 did not occur in isolation; the same advisory cycle also flagged an unrelated missing-encryption flaw in Apache Tomcat cluster communications and a pair of authentication-bypass vulnerabilities in N-able's N-central remote monitoring platform, underscoring that federal and enterprise patch teams were managing several unrelated, simultaneously exploited product families that week [1]. For organizations without a Federal Civilian Executive Branch remediation deadline, the KEV listing nonetheless functions as a reliable, independently corroborated signal that exploitation has moved from proof-of-concept demonstration to real-world attacker use, and should be treated with the same urgency internally regardless of whether a specific compliance deadline applies. The fact that meaningful exploitation activity preceded the public disclosure and KEV listing by roughly a month is itself a notable finding: organizations that wait for a KEV entry before beginning triage on newly disclosed AI infrastructure vulnerabilities may already be responding to an incident rather than preventing one.

Recommendations

Immediate Actions

Organizations operating Langflow OSS at version 1.10.0 or earlier should upgrade to 1.10.1 or later without delay, verifying the deployed version directly rather than trusting release notes or container tags, since prior Langflow patches have shipped with incomplete fixes that left the underlying flaw exploitable [3][7]. Any instance reachable from the public internet should be treated as a priority even if patching is already scheduled; where an immediate upgrade is not feasible, the instance should be removed from direct internet exposure and placed behind an authenticating reverse proxy or VPN, and the `LANGFLOW_AUTO_LOGIN` setting should be explicitly set to false so the platform requires real credentials rather than issuing tokens automatically [2][3]. Teams should also confirm, after patching, that the `/api/v1/auto_login` endpoint no longer issues tokens to unauthenticated callers, since verifying the actual runtime behavior is more reliable than trusting the version string alone [4].

Short-Term Mitigations

Given that exploitation attempts were recorded well before the July 17 disclosure, any Langflow instance that has been internet-facing since early July 2026 should be treated as potentially compromised pending log review. Security teams should audit access logs for unauthenticated calls to `/api/v1/auto_login` followed by POST requests to `/api/v1/validate/code`, review hosts for unauthorized cron entries, modified SSH `authorized_keys` files, or unexpected outbound network connections, and rotate every credential accessible to the Langflow instance, including LLM provider API keys, database credentials, and any cloud service account keys referenced in workflow configurations [3][6]. Organizations that rely on managed service providers or third-party integrators for Langflow-based AI pipelines should request confirmation of patch status and credential rotation from those providers directly, given the documented risk of cross-customer blast radius in multi-tenant deployments [1][7].

Strategic Considerations

The recurrence of unauthenticated RCE and authorization-bypass vulnerabilities in Langflow across four separate disclosures in 2026 indicates that organizations should not treat AI development and orchestration tooling as inherently lower-risk than production application infrastructure simply because it originated as a rapid-prototyping tool [6][7]. Security and platform teams should incorporate AI orchestration platforms into the same vulnerability management, network segmentation, and secrets

management programs applied to production systems, including enforcing least-privilege credential scoping so that a compromised orchestration instance cannot expose broader cloud or data resources than the specific workflow requires. Given the pattern of secrets aggregation inside these platforms, teams should also evaluate short-lived, vault-issued credentials over long-lived static keys embedded directly in workflow configuration, and should build recurring authorization and code-execution testing into the deployment lifecycle of any self-hosted AI agent platform rather than relying solely on vendor patch cadence.

CSA Resource Alignment

This incident extends a pattern CSA has tracked closely across Langflow's 2026 disclosure history, and two prior CSA research notes provide directly applicable context. [CVE-2026-33017: Langflow RCE Exploits Enterprise AI Pipelines](#), CSA's analysis of Langflow's first major unauthenticated RCE disclosure in 2026, documented the same downstream harvesting of LLM provider keys and cloud credentials that this incident's telemetry suggests is again occurring, and its mapping of the vulnerability class to MAESTRO's Layer 3 (Agent Frameworks and Orchestration) and Layer 4 (Deployment and Infrastructure) threat categories remains applicable to this newer chain. [Langflow Path Traversal: Unauthenticated RCE Actively Exploited](#), CSA's review of the June 2026 file-upload disclosure, traced attacker access back to the same root pattern evident here: default configurations and endpoint logic that assume network-level trust rather than enforcing authentication at the API layer, in this case the `/api/v1/auto_login` endpoint's failure to require credentials before minting a privileged session token. Taken together with the July 2026 CVE-2026-55255 cross-tenant authorization bypass that CISA also added to its KEV catalog that month [8], these four disclosures within a single year indicate the underlying authentication and code-execution model in Langflow has required repeated, piecemeal remediation rather than a single structural fix.

More broadly, CVE-2026-9198 illustrates a control gap squarely within the identity and access management and application security domains of CSA's AI Controls Matrix (AICM) v1.1: an endpoint intended purely for local development convenience was left reachable and privilege-granting in production-equivalent deployments, and a code-validation feature performed unsafe dynamic evaluation of untrusted input [10]. Organizations assessing their own AI orchestration platforms against AICM's identity, access management, and application and interface security domains should specifically verify that development-convenience features, such as auto-login or debug endpoints, are disabled or access-restricted before any instance is exposed to a shared or production network, and that any code-execution or validation feature enforces sandboxing rather than direct interpreter evaluation of user input.

References

- [1] The Hacker News. "[CISA Flags Langflow RCE, Tomcat, and N-central Flaws as Actively Exploited](#)." The Hacker News, August 2026.
- [2] SentinelOne. "[CVE-2026-9198: Langflow RCE Vulnerability](#)." SentinelOne Vulnerability Database, 2026.
- [3] Indusface. "[CVE-2026-9198: Critical Langflow RCE Under Active Exploitation](#)." Indusface Blog, 2026.
- [4] Field Effect. "[Langflow Vulnerability Chain Under Active Exploitation](#)." Field Effect Blog, 2026.
- [5] KEVIntel. "[CVE-2026-9198 Exploitation Observed – Langflow OSS](#)." KEVIntel, 2026.
- [6] Cloud Security Alliance. "[CVE-2026-33017: Langflow RCE Exploits Enterprise AI Pipelines](#)." Cloud Security Alliance, July 2026.
- [7] Cloud Security Alliance. "[Langflow Path Traversal: Unauthenticated RCE Actively Exploited](#)." Cloud Security Alliance, June 2026.
- [8] Qualys ThreatPROTECT. "[CISA Warns About Langflow Authorization Bypass Vulnerability Exploitation \(CVE-2026-55255\)](#)." Qualys ThreatPROTECT, July 2026.
- [9] Mallory.ai. "[Public PoC Released for Critical Langflow RCE Enabling SUPERUSER Access](#)." Mallory.ai, 2026.
- [10] Cloud Security Alliance. "[AI Controls Matrix \(AICM\) v1.1](#)." Cloud Security Alliance, 2025.
- [11] Langflow. "[Release v1.11.2](#)." GitHub, August 2026.