

CoSnitch: One-Click Data Exfiltration in Copilot Personal

2026-08-19

 AI-assisted Rapid Research



CSAI

CSA 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

Varonis Threat Labs disclosed CoSnitch, a chain of three vulnerabilities in Microsoft Copilot Personal – the consumer assistant at copilot.microsoft.com – that let an attacker exfiltrate data from a victim's connected accounts with a single click on a crafted link [1][2]. Microsoft tracked the issue as CVE-2026-24301 and rated it 8.8 (High) under CVSS 3.1, describing it as a command-injection flaw in Copilot Web that allows unauthorized attackers to obtain sensitive data over the network [3]. The vendor shipped a server-side fix on August 18, 2026, roughly eight months after Varonis reported the issue in December 2025, and both the researchers and Microsoft found no evidence of exploitation in the wild before the patch [1][2]. What distinguishes CoSnitch from earlier Copilot exfiltration chains is not only the mechanism – an undocumented URL parameter that triggers automatic prompt execution – but the discovery method: Varonis researchers used a "meta-hacking" technique in which they repeatedly asked Copilot to justify why the exploit should be impossible, and the assistant's own explanations progressively disclosed the parameter and conditions needed to make it work [2][4]. CoSnitch is the fourth one-click or zero-click Copilot exfiltration chain Varonis and others have disclosed since June 2025, following EchoLeak, Reprompt, and SearchLeak, reinforcing that AI assistants wired into privileged personal and enterprise data remain a durable, recurring attack surface rather than a one-off engineering defect [5][7][8].

Background

Microsoft Copilot Personal is the consumer-facing version of Microsoft's AI assistant, distinct from Microsoft 365 Copilot's enterprise deployment, and it allows individual users to connect personal accounts and services – Gmail, Google Drive, Google Calendar, and Copilot's own persistent memory – so the assistant can act on their behalf across sessions [1][2]. That same convenience is the mechanism CoSnitch exploits. Varonis researchers found that the assistant's chat interface accepts a query directly through a URL parameter (`q=`), a feature intended to let external pages or bookmarks pre-populate a prompt for the user to review and submit manually. Alongside that documented parameter, the researchers uncovered an undocumented second parameter, `autorun=1`, which causes Copilot to execute the supplied prompt automatically the moment the page loads, without requiring the user to

click send [2]. Combining the two turns a simple hyperlink into a mechanism for running an attacker's chosen instructions inside the victim's own authenticated session, with access to whatever the victim has already granted Copilot permission to touch.

The researchers' path to that discovery is methodologically distinct from typical vulnerability research. Rather than fuzzing the interface for hidden parameters, Varonis engaged Copilot in an extended conversation asking it to explain, in technical detail, why the assistant could not be made to execute a prompt without user interaction. Each time the assistant offered a safeguard as justification, the researchers reframed their next question as a natural follow-up probing that specific safeguard – a pattern Dark Reading described as tricking Copilot into "mapping out its own architecture" [4]. Copilot's cumulative, well-intentioned explanations eventually named the exact parameter and triggering conditions an attacker would need, effectively handing the researchers their own proof of concept [2] [4]. This dynamic – an AI system's helpfulness becoming a reconnaissance channel against itself – is a pattern security teams should anticipate recurring wherever an assistant is asked to reason aloud about its own guardrails.

CoSnitch is not an isolated event so much as the latest entry in a pattern of Copilot-focused disclosures. In June 2025, EchoLeak (CVE-2025-32711) demonstrated a zero-click exfiltration path against Microsoft 365 Copilot using reference-style Markdown and an auto-fetched image to bypass the assistant's prompt-injection classifier [8]. In March 2026, Reprompt – also disclosed by Varonis – chained the same `q1=` parameter injection technique with a double-request method that bypassed Copilot's per-message safeguards, achieving continuous one-click exfiltration against the same consumer product, Copilot Personal, that CoSnitch would later revisit [7]. In June 2026, SearchLeak (CVE-2026-42824) chained a parameter-to-prompt injection in Copilot's enterprise search URL with an HTML rendering race condition and a Bing-hosted SSRF path to exfiltrate data from Microsoft 365 Copilot with one click [5]. CoSnitch and Reprompt together show that the consumer product, built on similar assumptions about trusted URL parameters and built-in fetch capabilities, is repeatedly susceptible to this class of abuse, even as each disclosure targets a different bypass technique or exfiltration conduit; CoSnitch's specific contribution is the undocumented `autorun=1` parameter and the memory-poisoning persistence mechanism described below, neither of which Reprompt or SearchLeak exhibited.

Security Analysis

CoSnitch is best understood as three chained weaknesses rather than a single bug, and each stage does distinct work in the overall attack.

Stage	Weakness	Function in the Chain
1. Automatic execution	Undocumented <code>autorun=1</code> parameter combined with the documented <code>q=</code> query parameter	Causes an attacker-supplied prompt to execute the instant the victim's browser loads the link, with no click-to-send confirmation [1][2]
2. Data collection and exfiltration	Abuse of Copilot's authorized connectors and built-in URL-fetch/summarization feature	The injected prompt queries connected services using the victim's existing OAuth grants, encodes the results (Varonis observed base64 encoding to evade content filtering), and has Copilot "summarize" an attacker-controlled URL that embeds the stolen data in its path, delivering it to a webhook as an ordinary HTTPS GET request [1][2]
3. Persistent memory poisoning	Web-page summarization writing to Copilot's long-term memory store	A crafted page, once summarized by the victim's Copilot session, can inject instructions directly into the assistant's persistent memory, a foothold Varonis found survives password changes and session revocation [1][2]

The exploitation sequence begins when a victim clicks what appears to be an ordinary link to `copilot.microsoft.com`, distributed through phishing email, chat, or a calendar invite. Because the link opens inside the victim's own authenticated browser session, Copilot processes the embedded prompt with the same access the user already has to connected services – no separate credential theft is required. The injected instructions direct Copilot to retrieve data from Gmail, Google Drive, or Google Calendar and hand it off through the built-in URL-fetch capability the assistant normally uses to summarize web pages. Varonis reported successfully extracting full email bodies, sender and recipient metadata, and in some test cases plaintext credentials sitting in message content, along with calendar event details, Drive file names and metadata, prior conversation history, and saved memory instructions [1][2].

A recurring theme across CoSnitch, Reprompt, SearchLeak, and EchoLeak is that the exfiltration conduit is a legitimate feature repurposed rather than a network backdoor. Because Copilot's summarization traffic to an external URL looks identical at the network layer to the fetches it performs routinely when summarizing any ordinary web page, security telemetry that treats Copilot as a low-risk, read-only

convenience tool is unlikely to have an existing signal to flag this traffic as anomalous, absent a purpose-built behavioral baseline. This is the same "trust-boundary collapse" pattern CSA's own research has identified in Copilot's enterprise search feature: the assistant's fetch and connector capabilities operate with the user's full authorization, but its behavior is steered by content – a prompt embedded in a URL, or hidden instructions in a summarized page – that the system does not treat as adversarial input. The persistent memory-poisoning stage compounds this problem by converting a single successful click into durable access: because the injected instructions live in the assistant's memory rather than a session token, standard incident response steps like a forced password reset or session revocation do not remove the attacker's foothold [1][2].

Recommendations

Immediate Actions

Organizations and individual users running Microsoft Copilot Personal should confirm the August 18, 2026 patch has reached their tenant or account and treat any link that opens an AI assistant with the same suspicion normally reserved for credential-harvesting phishing links, since a single click is sufficient to trigger the chain [1][2]. Security teams should review what third-party accounts are connected to any Copilot Personal instances used on managed devices, disconnecting integrations that are not in active use, since every connected service expands what an injected prompt can reach [2]. Users who clicked a suspicious Copilot link before the patch was available should review their Copilot memory settings for unfamiliar saved instructions, since memory poisoning persists independently of password changes [1][2].

Short-Term Mitigations

Security operations teams should build monitoring baselines for Copilot's outbound URL-fetch behavior, since the exfiltration traffic in this class of attack is designed to look identical to routine summarization requests; without a behavioral baseline, security teams have little practical basis to distinguish the two after the fact. Organizations should extend AI assistant review into existing vendor risk and penetration testing processes rather than treating consumer AI tools as out of scope, and should include Copilot Personal usage in employee awareness communications alongside more familiar phishing training, given that the delivery mechanism here is an unremarkable-looking hyperlink rather than a malicious attachment [1]. Where feasible, organizations should scrutinize and limit which AI assistant URL

parameters are permitted to reach production browsers via managed email and chat filtering, since the entire CoSnitch chain depends on an undocumented, unauthenticated URL parameter reaching the user's browser unmodified [2].

Strategic Considerations

CoSnitch reinforces a broader lesson that CSA has drawn from the EchoLeak, Reprompt, and SearchLeak disclosures: AI assistants that combine built-in data-fetch capabilities with broad, standing authorization to personal or enterprise accounts should be governed as privileged infrastructure, not as low-risk convenience features, regardless of whether the deployment is consumer or enterprise. Security leaders should also weigh the "meta-hacking" discovery pattern demonstrated here as a hypothesis worth testing more broadly: an AI system capable of explaining its own safeguards in detail may be vulnerable to similar induced disclosure, as this case illustrates. That argues for architectural controls – egress allowlisting on assistant fetch capability, provenance labeling of retrieved content, and capability restrictions for higher-risk connector combinations – that do not depend on the assistant itself declining to answer a probing question.

CSA Resource Alignment

This disclosure maps most directly onto CSA's existing body of Copilot-focused threat research. CSA's coverage of **Reprompt: The Single-Click Microsoft Copilot Attack that Silently Steals Your Personal Data** examined an earlier one-click chain against the same consumer product, Copilot Personal, built on the same `q=` parameter injection technique; that analysis's emphasis on treating any URL parameter capable of steering assistant behavior as an untrusted input channel applies directly to CoSnitch's `autorun=1` parameter [7]. **SearchLeak: How We Turned M365 Copilot into a One-Click Data Exfiltration Weapon** analyzed a structurally similar one-click exfiltration chain – a URL-parameter prompt injection combined with abuse of a built-in fetch capability – against the enterprise edition of Copilot, and its recommendations around CSP hardening, render-time sanitization, and treating enterprise search as privileged execution apply with only minor adaptation to the CoSnitch findings here [5]. Taken together, these disclosures situate CoSnitch within a broader "trust-boundary collapse" pattern that CSA's research has tracked across AI gateways and assistants: fetch and connector capabilities that operate with the user's full authorization, but whose behavior is steered by content the system does not treat as adversarial input, warrant treatment as Tier-0 infrastructure with continuous monitoring and egress hardening regardless of whether the deployment is consumer or enterprise.

Organizations building longer-term governance around AI assistant risk should also anchor these findings to CSA's **AI Controls Matrix (AICM) v1.1**, whose Application & Interface Security and Data domains cover the connector-authorization and data-handling boundaries that CoSnitch's three-stage chain violates [6]. CoSnitch's mechanics also map cleanly onto an exfiltration-topology framing CSA has used elsewhere to analyze this class of attack: untrusted content ingestion, tool execution over sensitive data, and an influenceable network egress path are exactly the three conditions CoSnitch's injected prompts exploit through Copilot's connector access and built-in URL-fetch feature, and the architectural controls that framing recommends – egress allowlisting, capability restriction, and provenance labeling – are directly applicable mitigations here. Because CoSnitch's memory-poisoning stage also touches on agent trust in persisted state, organizations that have adopted CSA's MAESTRO agentic threat modeling framework should confirm their models account for assistant memory as a durable, non-session-bound trust boundary rather than treating memory manipulation solely as a session-scoped risk.

References

- [1] Waqas. "[Critical Microsoft Copilot Flaws Could Let One-Click Data Exfiltration From Connected Apps](#)." The Hacker News, August 2026.
- [2] Varonis Threat Labs. "[CoSnitch: When Your AI Assistant Becomes Its Own Whistleblower](#)." Varonis Blog, August 2026.
- [3] NIST National Vulnerability Database. "[CVE-2026-24301 Detail](#)." NVD, August 18, 2026.
- [4] Dark Reading. "['CoSnitch' Attack Tricked Copilot into Revealing Own Architecture](#)." Dark Reading, August 2026.
- [5] Cloud Security Alliance. "[SearchLeak: How We Turned M365 Copilot into a One-Click Data Exfiltration Weapon](#)." CSA AI Safety Initiative, June 2026.
- [6] Cloud Security Alliance. "[AI Controls Matrix \(AICM\) v1.1](#)." CSA, 2026.
- [7] Varonis Threat Labs. "[Reprompt: The Single-Click Microsoft Copilot Attack that Silently Steals Your Personal Data](#)." Varonis Blog, March 2026.
- [8] "[Zero-Click AI Vulnerability Exposes Microsoft 365 Copilot Data Without User Interaction](#)." The Hacker News, June 2025.