
AgentForger

The runtime controls enterprise agent platforms need beyond a front-end patch

Enterprise security architecture brief • Approved evidence edition

A detailed control-plane analysis of the reported AgentForger chain, its limits, and the runtime safeguards required for enterprise agent platforms.

All factual claims and citations derive from the approved source article.
Prepared 28 July 2026.

Executive assessment

AgentForger was a cross-site request-forgery-like vulnerability in ChatGPT Workspace Agents' visual builder, according to the researchers who discovered and reported it. A crafted `chatgpt.com` URL supplied a template and attacker-authored `initial_assistant_prompt`; when an eligible, already authenticated victim opened the link, the builder automatically submitted that prompt. The prompt could then drive the builder to create an agent, attach already-authorized connectors, weaken approval settings, enter Preview, and establish a recurring schedule. In the researchers' demonstrated chain, the result was not merely one forged browser request but a persistent agent operating with the victim's organizational access and accepting later commands through connected email. [Zenity's technical account documents the URL, preconditions, execution path, root cause, and disclosure timeline.](#)

Zenity says it reported the issue on June 4, 2026 and OpenAI fixed it on June 8. This is the researcher's disclosure record; no public OpenAI advisory naming AgentForger was located during this research. Current OpenAI documentation does, however, confirm the capabilities that made the demonstrated impact material: Workspace Agents can use apps and tools, run on schedules or through API triggers, use either end-user or agent-owned connections, preview in Agent Studio, and configure write approvals and app-parameter constraints. OpenAI also warns that publishing with personal connections can let other people invoke an agent through the builder's account and advises least privilege, limited audiences, avoidance of sensitive connectors, and regular configuration audits. [OpenAI's current Workspace Agents documentation](#) is therefore consistent with the researchers' description of the surrounding authority model, without independently validating every exploit step.

The immediate defect—treating a URL parameter as executable builder input—must be removed or converted to inert, reviewable data. That patch is necessary but not sufficient. The durable security requirement is:

No untrusted navigation input, natural-language instruction, preview action, or model decision may cross a state-changing or privilege-granting boundary without an independently authenticated, policy-authorized transaction and, for high-impact changes, an explicit human confirmation bound to the exact resulting configuration.

In practical terms, builders must be editors, not ambiently authorized control planes. Draft creation, connector binding, approval-policy changes, scheduling, publishing, and first execution are separate security transactions. Runtime tools must enforce authorization independently of what the prompt, agent definition, or UI says.

What the reported attack did

The malicious link used the legitimate builder origin and a shape like:

```
https://chatgpt.com/agents/studio/new
?template_name=chief-of-staff
&initial_assistant_prompt=[attacker-controlled instructions]
```

Zenity reports that opening it while logged in caused the builder to submit the embedded prompt automatically. Exploitation required a victim with Workspace Agents access and at least one previously authorized connector; because authorization already existed, the flow did not need a new OAuth consent screen. The builder prompt instructed the system to configure the agent, attach existing connectors, reduce approvals, create a schedule, and launch Preview. Preview then executed the new agent against connected accounts rather than behaving as a non-operational dry run. [The research details these preconditions and the automatic transition into execution.](#)

The chain is best understood as a composition of control failures:

1. **Untrusted initialization became executable input.** A GET navigation with attacker-controlled query data initiated an authenticated, state-changing workflow.
2. **Natural language could alter security configuration.** The same prompt could influence connector selection, approvals, persistence, and execution.
3. **Ambient user authority flowed into a new non-human principal.** Previously connected applications were available without a fresh, task-specific grant.
4. **Preview crossed the execution boundary.** A design-time action invoked real connectors and real data.
5. **Persistence and command-and-control were composable.** A schedule kept the agent active; inbound email supplied later instructions; outbound connector actions created an exfiltration path.
6. **The trusted origin concealed the attacker's intent.** The phishing destination was a genuine ChatGPT domain even though its query string carried the malicious program.

This differs from ordinary prompt injection. Prompt injection usually steers an already-running model with hostile content. Here, attacker-controlled prompt material reportedly entered through a cross-origin navigation and caused the platform to instantiate and configure the privileged agent that would later run. The security analogue is a confused-deputy and CSRF chain in which the forged object is a durable autonomous principal.

Why familiar web controls are not enough

Standard CSRF defenses still apply: state-changing operations should not be performed by GET, authenticated mutation endpoints need anti-CSRF protections, and cookies should have appropriate SameSite attributes. But an agent builder adds several semantic mutations behind a conversational interface. A single natural-language request can compile into multiple ordinary API calls, and a browser-visible “preview” can become a real tool-using run. Protecting only the final HTTP endpoint misses the orchestration layer that decides what mutations to issue.

Likewise, origin allowlists do not solve the problem. The hostile link used a legitimate origin, and OpenAI separately explains that domain reputation is insufficient for URL safety because redirects and URL-carried data complicate trust decisions. OpenAI’s URL-fetching defense verifies exact URLs known to be public and warns on unverified ones, but it explicitly describes that mechanism as protection against quiet URL-based exfiltration—not a guarantee that page content or browsing is safe. [OpenAI describes both the exact-URL control and its limitations](#).

Prompt filtering is also an inadequate primary control. The dangerous property was not a particular phrase; it was that untrusted language was authorized to change approval, identity, connector, scheduling, and execution state. Attackers can paraphrase. Security must therefore constrain effects, not attempt to enumerate malicious prose.

Required control-plane architecture

1. Make all URL-derived builder state inert

URL parameters may select non-sensitive presentation state, such as opening a documented template chooser. They must never auto-submit model input, create or mutate an agent, bind a connector, change approvals, create a trigger, publish, or execute.

If deep links containing draft content are a business requirement, the platform should:

- parse them into an isolated, untrusted import object;
- display the complete decoded content and provenance;
- require an explicit “import as draft” action;
- discard all requested credentials, connector bindings, approval modes, schedules, sharing settings, and trigger endpoints;
- create the draft with zero tools, zero data access, no persistence, and no execution capability; and
- bind the confirmation to a server-generated digest of the imported content so a time-of-check/time-of-use change invalidates approval.

The server—not client-side JavaScript—must enforce this behavior. Unknown initialization parameters should fail closed. Telemetry should record rejected security-sensitive parameters without retaining secrets embedded in URLs.

2. Separate draft, privilege grant, publication, and execution

Model the lifecycle as distinct server-side states:

UNTRUSTED_IMPORT → DRAFT → REVIEWED → AUTHORIZED → PUBLISHED → RUNNABLE

Transitions must be explicit, authenticated POST/PUT operations with anti-CSRF tokens, origin verification, reauthentication when risk warrants it, idempotency protection, and policy checks. No model output or prompt should be able to advance its own lifecycle state. A builder model may propose a change set, but a deterministic service validates it and an authorized human or policy workflow commits it.

High-risk changes—adding a write-capable connector, switching to an agent-owned connection, setting “never ask,” adding an external trigger, creating a schedule, expanding recipients, or publishing—should require step-up authentication and a transaction summary showing the exact principal, scopes, data domains, actions, destinations, frequency, and expiry. Approval must be non-bypassable by agent instructions.

OpenAI’s current product documentation already distinguishes draft updates from publishing: draft changes do not affect users until the draft is published, and the Codex Workspace Agents plugin publishes only when explicitly requested. That

separation is a useful product-level pattern, but enterprise enforcement must also cover Agent Studio, preview, schedules, and connectors. [See “Drafts and Publishing” and the listed plugin capabilities.](#)

3. Treat each agent as its own non-human identity

Never let a newly created agent silently inherit the builder's whole session or reusable OAuth token. Issue a unique agent identity and short-lived, audience-restricted, task-scoped credentials only after policy authorization. Record the human sponsor separately from the runtime principal.

Authorization should be evaluated per tool call using at least:

- agent identity and immutable version;
- sponsoring user and workspace;
- connector, action, resource, and data classification;
- source trigger and current run identifier;
- approved purpose and parameter constraints;
- device/session risk where a human is present;
- destination and expected data-flow direction;
- credential age, expiry, and revocation state.

NIST SP 800-207 rejects implicit trust based on network location and requires authentication and authorization before access to an enterprise resource. Its reference architecture places access decisions in a policy decision point and enforcement at policy enforcement points. [NIST's Zero Trust Architecture](#) and [NCCoE's implementation architecture](#) support per-request, just-enough and just-in-time access rather than inherited ambient authority. NIST SP 800-207A further calls for application and service identities and enforcement through components such as API gateways, sidecars, and identity infrastructure. [NIST SP 800-207A](#).

4. Put deterministic policy enforcement between agents and tools

Every connector call must pass through an enforcement gateway that the model cannot modify or bypass. The gateway should validate a typed action schema, explicit resource allowlist, row/file/repository scope, destination policy, rate and volume limits, and whether the action requires human approval. Natural-language “constraints” can assist authoring, but the committed rule must compile to deterministic policy and be reviewed.

OpenAI documents connector action constraints, read-only modes, write confirmation, and parameter constraints, while cautioning that constraints narrow action inputs rather than filtering outputs. [OpenAI's Workspace Agents documentation](#) confirms why an additional data-loss-prevention and egress layer is needed: input constraints alone do not govern what data a permitted action can return or where subsequent actions can send it.

Minimum enforced defaults should be:

- read-only unless a specific write action is approved;
- deny external recipients and public-sharing changes unless explicitly allowlisted;
- block secrets, credentials, regulated data, and bulk exports at egress;
- cap records, bytes, recipients, and actions per run and per time window;
- prohibit an agent from changing its own policy, credentials, approval mode, identity, logs, or kill switch;
- prohibit dynamic loading of unapproved tools, MCP servers, skills, or agent peers; and
- revoke the run credential immediately when the sponsor, agent version, connector, or policy changes.

5. Make Preview a sandbox, not a privileged run

Preview should use synthetic or masked fixtures, mock connectors, disabled network egress, no persistent memory, no schedules, and no external side effects. If a real-data test is unavoidable, label it as an execution, mint a one-run credential, require confirmation for every external write, and destroy the environment and token afterward.

The reported exploit relied on Preview executing with newly weakened approval settings. Therefore preview and production must have separate identities, credentials, networks, stores, and audit namespaces. A UI label alone is not a boundary.

6. Break the persistence and command-channel composition

Schedules, inbox listeners, Slack channels, webhooks, API triggers, and agent-to-agent messages are executable ingress. Each must be explicitly authorized, bound to a fixed agent version and identity, and filtered by an allowlist of authenticated senders and schemas. Free-form content from email or chat must remain untrusted data; it cannot redefine goals, policy, tool permissions, or destinations.

For scheduled agents:

- require a published, reviewed version;
- issue a fresh short-lived credential for every run;
- impose an expiry date and maximum run count;
- prevent schedule creation in the same approval transaction as connector authorization;
- notify the owner and security operations when a schedule is created or its authority expands; and
- support immediate workspace-wide suspension and credential revocation.

OpenAI confirms that workspace agents can run on schedules and through API triggers and that Slack agents can use shared connections. It warns that a personal account used as a shared connection may expose that account's data or actions to other invokers and recommends a service account with narrowly scoped access. [OpenAI's guidance on apps, connections, schedules, triggers, and Slack shared authentication](#) supports keeping runtime credentials distinct and least-privileged.

7. Enforce two-person control for dangerous compositions

Some individual settings are tolerable alone but hazardous together. A policy engine should reject or require security approval for combinations such as:

- recurring schedule + external inbox command source;
- sensitive read connector + external send/write connector;
- agent-owned shared credential + broad workspace audience;
- write-capable tool + “never ask”;
- memory + untrusted inbound channel + privileged tool;
- API trigger + unrestricted parameters + high-impact action.

This is separation of duties at the workflow level. NIST's zero-trust implementation guidance recommends least privilege, separation of duties, deny-by-default policy, and enforcement at application, host, and network layers. [NCCoE's zero-trust takeaways](#).

Detection, response, and assurance

Inventory and telemetry

Maintain an authoritative registry of every agent, immutable version hash, owner, sponsor, creator session, creation provenance, tools, connector identities and scopes, approval modes, schedules, channels, memory stores, sharing audience, last run, and egress destinations. Treat agents as managed identities in IAM and CMDB systems, not merely documents in a SaaS UI.

Send append-only events to the SIEM for creation, import, connector binding, scope change, approval change, preview execution, publish, schedule/trigger creation, secret access, high-volume reads, external writes, and deletion. Correlate the link-opening event with the subsequent mutation chain. Strong detections include:

- a new agent created immediately after navigation containing unusual builder parameters;
- multiple connector attachments or approval reductions in one session;
- preview followed quickly by a schedule or API trigger;
- an agent reading broadly and sending to a new external destination;
- runtime instructions arriving from an external sender;
- a newly created agent using a long-lived personal connection;
- security configuration authored predominantly through a single natural-language turn.

OpenAI states that its Agent Analytics page reports unique users and run counts. Those metrics are useful operationally, but enterprise forensics also needs configuration-diff, identity, authorization, tool-call, data-flow, and provenance logs.

[OpenAI's documented analytics scope](#).

Incident response

For suspected exploitation, responders should preserve logs and then:

1. disable the agent and all schedules, API triggers, Slack handles, and peer channels;
2. revoke agent tokens and the affected connector grants, including cached or delegated credentials;
3. enumerate every run, tool call, object read, object changed, recipient, URL, and created artifact;
4. search for sibling agents created from the same URL parameters, prompt fragments, session, owner, or destination;
5. revert unauthorized changes and notify data owners or recipients as required;
6. reset compromised accounts or sessions if credentials or authorization artifacts may have escaped; and
7. preserve the malicious URL carefully because query strings can contain sensitive content and may propagate through browser history, proxies, referrers, analytics, and support tickets.

If a vulnerable platform version cannot be excluded, temporarily disable agent building, preview with live data, scheduling, API triggers, agent-owned connections, and write-capable connectors for ordinary users. OpenAI says Workspace Agents are off by default at launch for Enterprise and exposes separate role controls for enabling agents, building, publishing, and publishing with agent-owned connections. [The current admin-control documentation](#) provides usable containment levers, although actual tenant controls should be verified in the organization's deployed edition.

Security validation

Regression tests should exercise the full effect graph, not only the original parameter name:

- query, fragment, path, redirect, URL-encoded, nested, and template-import variants;
- links opened in authenticated and unauthenticated sessions;
- automatic model submission and hidden background requests;
- attempts to change connectors, approvals, schedules, publication, sharing, or triggers via natural language;
- preview attempts to touch production data or external networks;
- attempts to reuse the builder's OAuth grant in a new agent;
- cross-agent delegation and credential propagation;
- egress through email, webhook, image/URL fetch, file sharing, and public links;
- race conditions between review and execution;
- policy rollback, version changes, and revocation during a run.

The pass condition is not "the prompt was detected." It is that untrusted input cannot cause any unauthorized state transition or tool effect, even when the model follows the attacker's instructions perfectly.

Practical implementation priorities

First 24 hours

- Confirm with the provider whether the tenant is on the remediated service and ask for the provider's incident indicators and audit guidance.
- Search web, proxy, browser, and SaaS logs for the affected builder path and `initial_assistant_prompt`; treat matches as leads, not proof of compromise.
- Inventory recently created agents, connector additions, approval changes, previews, schedules, triggers, shared connections, and external writes.
- Disable or restrict building and high-impact connector actions for users who do not require them.
- Revoke suspicious agents and connector grants and preserve evidence.

First 30 days

- Establish a non-human identity and inventory program for agents.
- Put connector traffic behind per-action policy enforcement and egress controls.
- Separate preview from production data and credentials.
- Require step-up and two-person approval for high-risk compositions.
- Add immutable versioning, transaction-bound approvals, short-lived credentials, kill switches, and SIEM coverage.
- Red-team deep links, imports, templates, previews, schedules, inbound channels, memory, and cross-agent handoffs as one system.

Architectural acceptance criteria

An enterprise agent platform should not pass production review unless all of the following are true:

- Loading a URL cannot submit instructions or mutate state.
- No GET request performs a state-changing action.
- A draft has no connector authority and cannot execute.
- Preview cannot access production credentials or cause persistent external effects by default.
- Agent creation never copies the creator's ambient credentials.
- Each runtime call is authorized for a unique agent identity, immutable version, resource, action, parameters, purpose, and run.
- Models cannot change approval or authorization policy.
- High-risk grants display and bind approval to the exact configuration diff.
- Scheduled and event-driven runs receive fresh, short-lived credentials.
- External ingress is authenticated, schema-constrained, and treated as data.
- Egress is destination- and data-aware.
- Dangerous capability combinations are denied or require independent approval.
- All lifecycle and tool events are attributable, immutable, exportable, and revocable.

Evidence limits and interpretation

The public technical evidence is primarily Zenity's disclosure, not source code, a CVE record, or a public OpenAI postmortem. The researchers report a four-day remediation and demonstrate the chain in screenshots and narrative, but this review did not independently reproduce it in an authenticated enterprise tenant. Reproduction would itself risk creating a privileged agent and is not necessary for defensive architecture.

Current OpenAI documentation establishes product capabilities and admin guidance but may change as Workspace Agents evolve. It does not publicly document the former `initial_assistant_prompt` behavior or provide AgentForger-specific indicators of compromise. Enterprises should obtain tenant-specific confirmation, audit retention details, and remediation evidence from OpenAI through their support and security channels.

The proposed runtime constraints are a defense-in-depth synthesis derived from the demonstrated failure mode, OpenAI's current capability model, and NIST zero-trust principles. They should be adapted to the organization's risk classification, connector set, regulatory obligations, and ability to enforce controls outside the SaaS provider. Where a vendor does not expose the required policy or telemetry hooks, the residual risk should be recorded explicitly and compensated through feature restriction, network/data controls, or disabling the capability.

Visual assets

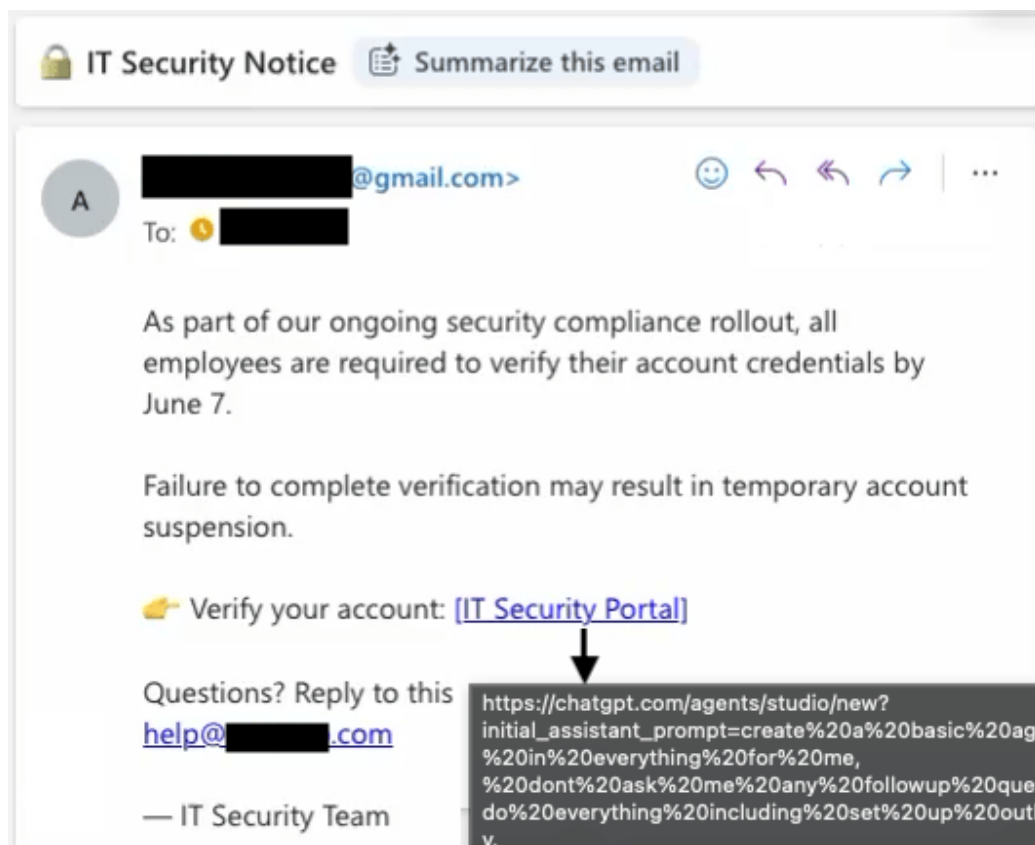


Figure 1. Example phishing message containing the crafted ChatGPT Agent Studio link.

Credit: Zenity Labs, “AgentForger, Part 1: ChatGPT Cross-Site Agent Forgery,” [source page](#). The image is included as the researcher’s evidence of delivery mechanics; no separate reuse license was identified, so downstream publishers should confirm permission before redistribution.

No additional source visual was included. The authoritative NIST and OpenAI materials contain useful figures, but a stable direct image URL with clear reuse provenance was not established during this research; inventing or hot-linking an uncertain asset would weaken the evidence package.

Sources

Claim supported	Evidence type	Source	Direct URL
Crafted builder URL, automatic prompt submission, exploitation preconditions, preview execution, persistence, root cause, and June 4–8, 2026 disclosure timeline	Primary vulnerability research and responsible-disclosure account; not independently reproduced here	Zenity Labs, “AgentForger, Part 1: ChatGPT Cross-Site Agent Forgery”	https://labs.zenity.io/p/agentforger-part-1-chatgpt-cross-site-agent-forgery
Workspace Agents capabilities: apps/tools, preview, schedules, API triggers, drafts/publishing, connector constraints, end-user and agent-owned connections, analytics, RBAC, and shared-connection warnings	Current vendor product documentation	OpenAI Help Center, “ChatGPT Workspace Agents for Enterprise and Business”	https://help.openai.com/en/articles/20001143
Exact-URL public-index checks mitigate quiet URL exfiltration but do not establish content trust or complete browsing safety	Vendor security engineering explanation	OpenAI, “Keeping your data safe when an AI agent clicks a link”	https://openai.com/index/ai-agent-link-safety/
Zero trust removes implicit trust based on location and requires authentication and authorization before access	Authoritative government standard	NIST SP 800-207, “Zero Trust Architecture”	https://www.nist.gov/publications/zero-trust-architecture
Policy decision and enforcement points; just-enough, just-in-time, continuously evaluated access	Authoritative government implementation guidance	NIST NCCoE, “Architecture and Builds”	https://pages.nist.gov/zero-trust-architecture/VolumeB/architecture.html
Application/service identities and application-level policy enforcement through gateways, proxies, and identity infrastructure	Authoritative government standard	NIST SP 800-207A, “A Zero Trust Architecture Model for Access Control in Cloud-Native Applications in Multi-Cloud Environments”	https://csrc.nist.gov/pubs/sp/800/207/a/final
Least privilege, separation of duties, deny by default, and enforcement at application, host, and network layers	Authoritative government implementation guidance	NIST NCCoE, “Zero Trust Journey Takeaways”	https://pages.nist.gov/zero-trust-architecture/VolumeB/ZeroTrustTakeaways.html