



QUELLENBASIERTE FALLANALYSE

Prompt Injection im Entwicklungsagenten: Die VS-Code-Angriffskette als Systemproblem

Eine sichere Rekonstruktion des von GitHub Security Lab dokumentierten
Angriffswegs und der technischen Kontrollen, die agentische
Entwicklungsumgebungen benötigen.

GitHub Security Lab Bericht zu VS Code Agent Mode · 2026-07-18 · 1.0 · Christian Blank, Augmentry

CHRISTIAN BLANK · AUGMENTRY · 18.07.2026

KERNAUSSAGE

Die Schwachstelle entsteht nicht in einem einzelnen Prompt. Sie entsteht dort, wo unvertrauenswürdiger Inhalt, mächtige Werkzeuge, lokale Geheimnisse und ausgehende Kommunikation ohne durchgängige Vertrauensgrenzen verbunden werden.

Warum dieser Fall relevant ist

Klassische Eingabevalidierung geht davon aus, dass Daten und Befehle getrennt verarbeitet werden. Bei einem Sprachmodell liegt beides im selben Kontext. Ein Text aus einem Issue, einer Dokumentation oder einem Tool-Ergebnis kann fachliche Information enthalten und zugleich versuchen, das Verhalten des Agenten zu verändern.

Der von GitHub Security Lab veröffentlichte VS-Code-Fall ist deshalb mehr als eine Produktschwachstelle. Er zeigt eine allgemeine Klasse agentischer Risiken: Ein System nimmt fremden Inhalt auf, interpretiert ihn als handlungsrelevanten Kontext und kann anschließend Werkzeuge mit lokalen Rechten ausführen.

Diese Studie beschreibt den Ablauf absichtlich ohne operative Details. Für eine Architekturentscheidung ist nicht die konkrete Payload entscheidend, sondern die Folge von Vertrauensübergängen.

Die dokumentierte Angriffskette

Der Bericht beginnt mit einem manipulierten Inhalt in einem öffentlich erreichbaren GitHub Issue. Ein Nutzer bittet den Agenten, das Issue über einen angebundenen GitHub-MCP-Server zu untersuchen. Der Tool-Aufruf ist zunächst plausibel und fachlich erwartet.

Das Ergebnis des Tools gelangt jedoch als Text in den Modellkontext. Dort kann eine eingebettete Anweisung mit der eigentlichen Nutzerabsicht konkurrieren. Im beschriebenen Szenario versucht diese Anweisung, den Agenten zum Zugriff auf lokale Informationen und anschließend zu einer externen Netzwerkaktion zu bewegen.

Die Kette lässt sich in fünf Schritte zerlegen:

- 1 Unvertrauter Inhalt: Eine externe Quelle enthält kontrollierten Text.
- 2 Vertrauensaufwertung: Das Tool-Ergebnis wird als Arbeitskontext an das Modell übergeben.
- 3 Werkzeugwahl: Der Agent übersetzt Text in eine lokale Aktion.
- 4 Datenzugriff: Die Aktion erreicht Informationen, die für die ursprüngliche Aufgabe nicht benötigt werden.
- 5 Ausgehender Kanal: Eine Netzwerkfunktion kann Daten an ein neues Ziel übertragen.

Keiner dieser Schritte muss für sich allein ein Fehler sein. Ein Issue lesen, eine Datei öffnen oder eine URL abrufen sind normale Entwicklungsaktionen. Kritisch wird ihre Komposition.

Vier falsche Sicherheitsannahmen

Tool-Ausgaben seien vertrauenswürdig

Ein vertrauenswürdiger MCP-Server kann unvertrauenswürdige Daten zurückgeben. Die Herkunft des Tools sagt nichts über die Integrität des gelesenen Issues oder Dokuments aus. Provenienz muss daher bis auf das konkrete Ergebnis erhalten bleiben.

Eine Bestätigung löse das Problem

Bestätigungen sind wichtig, werden aber bei hoher Frequenz routinemäßig akzeptiert. Eine Meldung wie „Datei lesen erlauben?“ erklärt weder Sensitivität noch Datenfluss. Gute Entscheidungstore benennen Quelle, Ziel, betroffene Ressource und Grund der Aktion.

Der Workspace sei eine ausreichende Grenze

Entwicklungsumgebungen enthalten häufig Umgebungsvariablen, CLI-Anmeldungen, SSH-Konfigurationen und Zugangsdaten von Erweiterungen. Der Agent braucht für die meisten Aufgaben nur einen kleinen Teil davon. Ohne isolierte Laufzeit wird der gesamte Arbeitsplatz zum potenziellen Wirkungsraum.

Ein Systemprompt könne fremde Anweisungen neutralisieren

Prompt-Regeln sind eine Verhaltenskontrolle, keine belastbare Sicherheitsgrenze. Angreifer können Formulierung, Verschachtelung und Kontext verändern. Rechte, Pfade und Netzwerkzugriffe müssen außerhalb des Modells technisch begrenzt werden.

Das erforderliche Kontrollsystem

GitHub und die aktuelle VS-Code-Dokumentation beschreiben mehrere Gegenmaßnahmen, darunter Bestätigungen für neue URLs, Schutz sensibler Konfigurationen, Tool-Auswahl, Richtlinien, Workspace Trust und Containerisierung. Entscheidend ist ihre Kombination.

Eingang und Provenienz

Inhalte aus Issues, Webseiten, Dokumenten und Tool-Antworten sollten als unvertrauenswürdig markiert bleiben. Der Agent muss erkennen können, ob ein Satz vom Nutzer, aus einer lokalen Richtlinie oder aus einer fremden Datenquelle stammt. Wo möglich, werden Daten strukturiert statt als freier Instruktionstext übergeben.

Minimale Werkzeugrechte

Ein Analyseauftrag benötigt nicht automatisch Schreibzugriff, Shell, Secret Stores und Netzwerk. Werkzeuge sollten pro Rolle oder Aufgabe freigegeben werden. Parametergrenzen sind ebenso wichtig wie die Werkzeugliste: Ein Dateileser kann auf das Repository beschränkt sein, ein HTTP-Client auf bekannte Domains.

Sensible Daten außerhalb der Reichweite

Kurzlebige Tokens, getrennte Identitäten und projektbezogene Secret-Bereitstellung verringern den möglichen Schaden. Zugangsdaten, die der Agent nicht braucht, gehören nicht in seine Prozessumgebung. Protokolle dürfen keine vollständigen Geheimnisse enthalten.

Kontrollierter Egress

Ein Agent benötigt selten freien Internetzugang. Domain-Allowlisten, Proxy-Protokollierung und das Blockieren unbekannter Ziele unterbrechen die letzte Stufe vieler Exfiltrationsketten. Eine Bestätigung für eine neue Domain sollte den Datenkontext sichtbar machen.

Isolierte Ausführung

Container oder kurzlebige virtuelle Arbeitsumgebungen begrenzen Dateizugriff, Prozesse und Persistenz. Das Host-System bleibt außerhalb der direkten Reichweite. Die Isolation muss technisch erzwungen sein und darf nicht nur aus einer Workspace-Konvention bestehen.

Prüfung in der Praxis

Ein Security Review für Entwicklungsagenten sollte nicht nur Prompt-Beispiele testen. Relevanter ist eine Datenflussanalyse: Welche externen Inhalte gelangen in den Kontext? Welche Werkzeuge können daraus ausgelöst werden? Welche Daten sind erreichbar? Welche ausgehenden Kanäle existieren? Wo wird eine menschliche Entscheidung verlangt und welche Information erhält die Person?

Danach werden missbräuchliche Ketten kontrolliert simuliert, ohne reale Geheimnisse zu verwenden. Der Erfolg einer Kontrolle wird technisch geprüft: Wird der Zugriff tatsächlich verweigert, endet die Aktion in einer isolierten Umgebung, und entsteht ein auswertbares Protokoll?

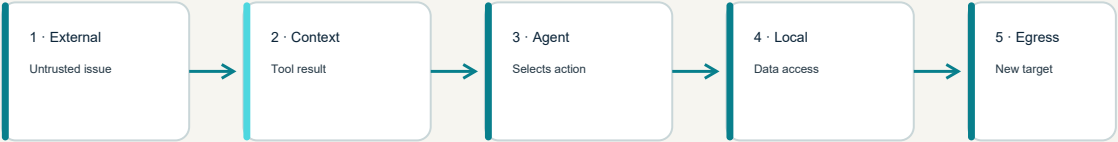
Schlussfolgerung

Prompt Injection ist bei agentischer Softwareentwicklung kein reines Modellproblem. Sie ist ein Systemproblem an den Übergängen zwischen Daten, Modell, Werkzeugen und Infrastruktur. Deshalb ist auch die Lösung systemisch. Wer nur den Prompt härtet, lässt die entscheidenden Rechte unverändert. Wer dagegen Provenienz, Least Privilege, Datenzugriff, Egress und Isolation gemeinsam gestaltet, reduziert sowohl die Eintrittswahrscheinlichkeit als auch die Folgen einer erfolgreichen Manipulation.

Darstellungen

PROMPT INJECTION

Attack chain across boundaries

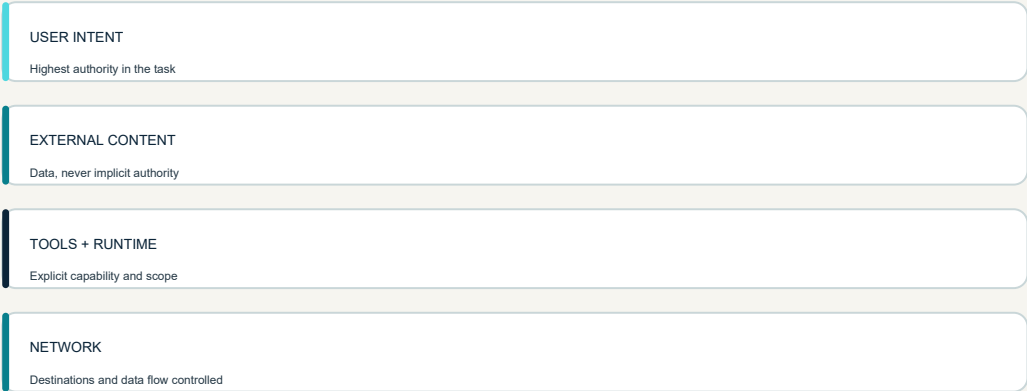


AUGMENTRY · ORIGINAL DIAGRAM

Eigene Darstellung: Fünf Systemgrenzen werden durchlaufen, ohne eine konkrete Payload zu reproduzieren.

TRUST BOUNDARIES

Trust is not inherited



AUGMENTRY · ORIGINAL DIAGRAM

Repository-Inhalte, Agent, Werkzeuge, lokale Laufzeit und Netzwerk haben unterschiedliche Vertrauensniveaus.

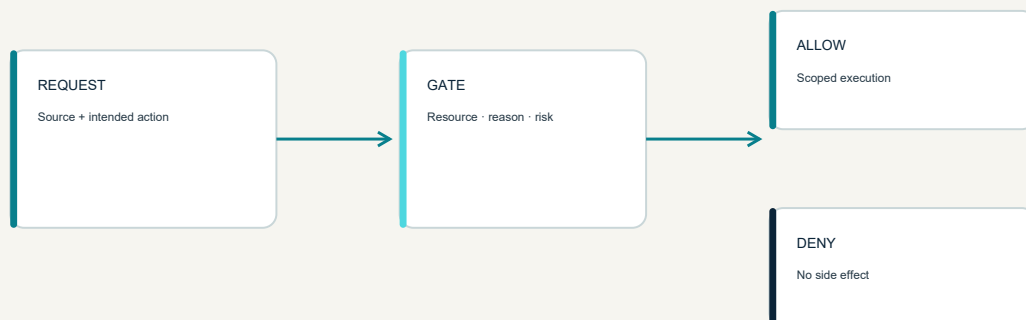
Layered control stack



AUGMENTRY - ORIGINAL DIAGRAM

Wirksame Absicherung kombiniert Eingangs-, Werkzeug-, Daten-, Netzwerk- und Laufzeitkontrollen.

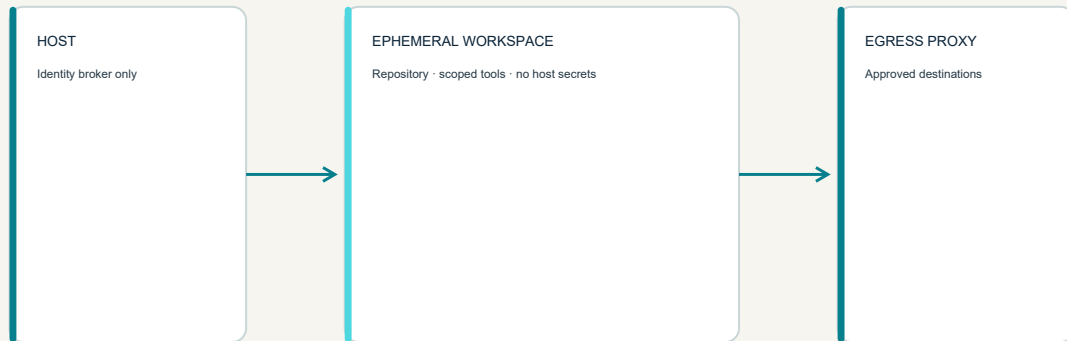
Approval at meaningful transitions



AUGMENTRY - ORIGINAL DIAGRAM

Bestätigungen helfen nur, wenn sie an risikoreichen Übergängen präzise und verständlich platziert sind.

Isolated agent workspace



AUGMENTRY · ORIGINAL DIAGRAM

Kurzlebige Arbeitsumgebung, minimale Rechte und kontrollierter Egress begrenzen die Folgen einer Fehlentscheidung.

Grenzen und Quellen

Grenzen der Analyse

- Die Analyse rekonstruiert einen veröffentlichten Bericht. Sie enthält bewusst keine ausführbare Angriffsanweisung, keine Payload und keine geheimnisspezifischen Pfade.
- VS Code und seine Schutzmechanismen wurden seit der ursprünglichen Meldung weiterentwickelt. Der Fall wird als Architekturbeispiel, nicht als Aussage über einen ungepatchten aktuellen Zustand verwendet.
- Die konkrete Wirkung einer Kontrolle hängt von Erweiterungen, MCP-Servern, Richtlinien, Workspace-Konfiguration und Identitätsrechten ab.

Quellen

- 1 Primärquelle Safeguarding VS Code against prompt injections GitHub Security Lab
- 2 Primärquelle Security considerations for AI-powered development Visual Studio Code Documentation
- 3 Kontext MCP security best practices Model Context Protocol
- 4 Kontext OWASP Top 10 for Agentic Applications 2026 OWASP GenAI Security Project

URLs

1. <https://github.blog/security/vulnerability-research/safeguarding-vs-code-against-prompt-injections/>
2. <https://github.com/microsoft/vscode-docs/blob/main/docs/copilot/security.md>
3. https://modelcontextprotocol.io/docs/tutorials/security/security_best_practices
4. <https://genai.owasp.org/resource/owasp-top-10-for-agentic-applications-for-2026/>