



Künstliche Intelligenz im Security-Kontext: Fluch oder Segen



Agenda

- Vorstellung und Einleitung
- KI als Werkzeug für Angreifer
- KI zur Verbesserung der Sicherheit
- Neue Probleme und Schwachstellen durch bzw. in KI
- Produkte zur Absicherung der neuen von KI verursachten Probleme
- Fazit

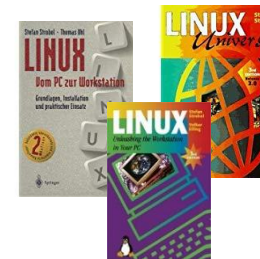


Über den Referenten



- IT seit 1983
 - Apple IIe, später erste PCs, Software-Entwicklung
- 1989
 - erste eigene Firma im Bereich Software-Entwicklung
- 1990-95
 - Studium Medizinische Informatik in Heidelberg / Heilbronn
 - Seit 1992 Internet, Unix, Linux
- 1993
 - erste Buchveröffentlichung „Linux vom PC zur Workstation“
 - Übersetzungen auf Englisch, Japanisch, Chinesisch etc.
- 1993/94
 - DESS Intelligence Artificielle am LIA in Chambéry / Frankreich
- 1995
 - Fokus auf IT-Sicherheit, erste Firewall-Installationen
 - Mitgründung der Centaur Communication GmbH, Fokus Internet / Sicherheit
- 1997
 - Veröffentlichung des ersten Firewall-Buches
 - Beginn der Nebentätigkeiten als Dozent an Hochschulen
- 2002
 - Mitgründung der cirosec GmbH, Fokus Informationssicherheit, Penetrationstests, Red Teaming, Incident Response, Beratung, Konzepte, Lösungen

Seit >40 Jahren IT



Seit >30 Jahren KI

Seit >30 Jahren Security



Seit 24 Jahren cirosec

Wer ist cirosec?

- Innovatives Unternehmen mit Fokus auf Informationssicherheit
 - Gegründet 2002 mit Sitz in Heilbronn, inzwischen ca. 80 Mitarbeiter
- Beratung, Konzepte, Reviews, Analysen
 - Architekturen, Zero Trust, Erkennung und SOC, sensible Daten, ...
- Red Teaming und Penetrationstests
 - In fast allen Bereichen
- Incident Response und Forensik
 - Von Konzepten bis SLA
- Auswahl und Implementierung von Produkten und Services
 - Konzepte, neutrale Evaluationen und Umsetzungen
- Trainings und Awareness
- Wir leben vom Know-how unserer Mitarbeiter
 - Erfahrene Spezialisten, Buchautoren

Veröffentlichungen und Vorträge

- Diverse Fachbücher
- Artikel vor allem iX und ct,
aber auch Computerwoche, kes, IT-Sicherheit usw.
 - Z. B.: iX-Sonderheft Security
- Viele Vorträge auf Konferenzen
 - Sowohl in Deutschland als auch im Ausland



Beispiele für Fachartikel der letzten Jahre

IX MAGAZIN FÜR PROFESSIONELLE INFORMATIONSTECHNIK

Analyse: Ein dreifacher Riegel für Hacker mit Microsofts neuem Defender

06.10.2020 10:43 Uhr
Hagen Molzer



(Bild: anrb/Shutterstock.com)

Im Rahmen der Ignite hat Microsoft Namen bietet der Defender aus

Auf seiner Ignite zeigte Microsoft Neustrukturierung sowie die W umfasst der Microsoft 365 Defen Dienste:

- Defender for Endpoint (vor
- Defender for Office 365 (v
- Defender for Identity (vom

Diese bilden zukünftig zusammen Center – den Umfang der XDR-S Sentinel, das als SIEM-Tool auch Gesamtbild integrieren und so d

Klassifikation: Intern

REPORT | CLOUD-SICHERHEIT

Cloud-Sicherheit auf dem Prüfstand

Joshua Tiago



Wolken ohne Regen

Die drei größten Cloud-Plattformen von Amazon, Microsoft und Google realisieren unterschiedliche Sicherheitskonzepte – und genau das ist für Unternehmen von großer Relevanz.

Ob im Unternehmensumfeld oder im Privathaushalt: Kaum eine App, ein IoT-Gerät oder auch Fahrzeug kommt ohne Cloud aus. Die drei größten Anbieter der Cloud-Lösungen sind Amazon (AWS), Microsoft (Azure) und Google (GCP). Die großen Anbieter und die ihnen alle gängigen Szenarien für Cloud-Umgebungen ab.

Während der Pandemie hat sich die Cloud für viele Arbeitnehmer als wesentliche Arbeitsmittel etabliert. Zudem lassen Unternehmen, die vom Homeoffice-Trend erfasst wurden, dank Cloud-Angeboten immer mehr in der beruflichen Alltag mitzudenken. Kein Wunder also, dass für das Jahr 2020 der weltweite Umsatz der großen Cloud-Anbieter auf 300 Mrd. US-Dollar geschätzt wurde. Bevor die Cloud-Anbieter wirklich waren, standen jedoch für viele Sicherheitsexperten in Unternehmen einige Fragen im Raum: Wie sicher sind meine Daten? Wie hat Zugriff auf Daten? Wo lassen Gefahren nach? Die Fragen sind zwar noch heute die große Unsicherheit ist, dass die

- TRAC**
- AWS, Azure und GCP sind die drei größten Cloud-Plattformen, die alle gängigen Szenarien für Cloud-Umgebungen abdecken.
 - Die Vergabe von Benutzerrollen kann in AWS verwirrend sein. Unternehmen sollten dafür zuerst ein Konzept entwickeln.
 - Ein Szenario und Gruppen sind zu verwalten, kann ein Administrator in der Azure-Umgebung das interne Active Directory mit dem Azure Active Directory koppeln.
 - Ein Security Command Center von GCP kann eine sicherheitsrelevante Informationen auf einen Blick vereinen.

PRAXIS | ADSICHERHEIT

Active Directory mit dem Schichtenmodell schützen

Von Hagen Molzer



Das Active Directory (AD) ist das zentrale Element der IT-Infrastruktur. Es ist die Basis für die Authentifizierung und Autorisierung aller Benutzer und Dienste im Netzwerk. Das neue Tutorial zeigt, wie man es umsetzt.

Erkennen und reagieren

Neue Verteidigungsansätze: EDR und XDR

Endpoint und Extended Detection and Response (EDR, XDR) sollen kompromittierte Systeme erkennen und die Incident Response unterstützen. Was leisten diese noch jungen Techniken wirklich?

Die Idee ist, dass die Hackerangriffe nur ein Teil der Bedrohungslage sind. Die meisten Angriffe werden durch Software-Schwachstellen ermöglicht, die in der Regel nicht erkannt werden. EDR und XDR zielen darauf ab, diese Schwachstellen zu identifizieren und zu beheben.

TRAC

- Ein Active Directory-Denkmal zu schützen, bedeutet es mehr als Absätze und Hardwaremaßnahmen. Microsoft hat ein Sicherheitskonzept entwickelt, das auf die Absicherung der gesamten IT-Infrastruktur abzielt.
- Grundlage des Konzepts ist die Festlegung eines zentralen Tiers, das die gesamte IT-Infrastruktur abdeckt und die verschiedenen Ebenen der IT-Infrastruktur abdeckt.
- Die Einführung des Tier-Modells führt zu verschiedenen Änderungen in der IT-Infrastruktur und ist sehr wichtig für die Absicherung und den Betrieb der IT-Infrastruktur.



TRAC

- AWS, Azure und GCP sind die drei größten Cloud-Plattformen, die alle gängigen Szenarien für Cloud-Umgebungen abdecken.
- Die Vergabe von Benutzerrollen kann in AWS verwirrend sein. Unternehmen sollten dafür zuerst ein Konzept entwickeln.
- Ein Szenario und Gruppen sind zu verwalten, kann ein Administrator in der Azure-Umgebung das interne Active Directory mit dem Azure Active Directory koppeln.
- Ein Security Command Center von GCP kann eine sicherheitsrelevante Informationen auf einen Blick vereinen.

REPORT | IT-SICHERHEIT

Doppeltes Spiel

Gefahren durch Angriffe auf und mit KI

Stefan Strobel

Wie vieles in der IT-Sicherheit hat auch der Einsatz künstlicher Intelligenz zwei Seiten. Zwar kann sie helfen, Angriffe automatisch zu erkennen und zu abwehren, aber sie kann auch selbst zum Ziel von Angriffen werden.



TITEL | ENDGERÄTESICHERHEIT

Auf dem Radar

Endpoint Detection and Response: Gefahren schnell erkennen und reagieren

Konstantin Bücheler, Martin Hartmann, Alain Rödel, Stefan Strobel

Aus den einst simplen Virenskannern, die heutigen Sicherheitsanforderungen nicht mehr genügen, sind komplexe Produkte zum Schutz der Endgeräte im Netz entstanden. IX hat sich den spiggen Markt der EDR-Produkte angeschaut.



TRAC

- Gegen heute gezielte und komplexe Angriffe reichen primäre Sicherheitsmaßnahmen nicht mehr aus. Je früher ein Unternehmen einen Angriff auf die eigene IT erkennt, desto besser ist der Schaden begrenzt.
- Altere Systeme wie IDS oder Virenschutz können nicht umgangen werden und werden durch moderne Sicherheitskonzepte ersetzt.
- Neuere Systeme wie EDR und XDR können und integrieren zahlreiche Daten, um eine umfassende Analyse der Bedrohungslage zu ermöglichen und die Bedrohungslage möglichst gering zu halten.

Prävention als Voraussetzung

Natürlich bedeutet das nicht, dass Produkte von Privilegien von Privilegien. Unternehmen sollten sichergehen, dass sie die richtigen Maßnahmen ergreifen, um die Bedrohungslage zu minimieren und die Bedrohungslage möglichst gering zu halten.

Management und Wissen | Zero-Trust-Implementierung

Vertrauen ist gut, kein Vertrauen ist besser?!

Entwicklung und Implementierungen von „Zero Trust“

Zero-Trust-Modelle verabschieden sich von der klassischen Unterscheidung vertrauenswürdiger Netze, Geräte oder Benutzer und hinterfragen mehr oder minder jeden einzelnen Zugriff. Unser Autor erläutert, in welchen Systemen und Lösungen dieser Ansatz heute zu finden ist.

Von Stefan Strobel, Heilbronn

Vertrauen ist eine wichtige Grundlage für die Zusammenarbeit in Unternehmen und auch in der Informationssicherheit ist es gewohnt, Personen, Systemen, Daten oder Netzwerken zu vertrauen. Beim aktuellen Trend „Zero Trust“ soll es aber zumindest kein implizites Vertrauen mehr geben. Um diese Tendenz (etwas anders) zu verstehen, lohnt es sich, die Entwicklung der letzten Jahre zu betrachten.

Man muss sich letztlich fragen, welche Teil der IT sich überhaupt noch „innerhalb“ einer Organisation befindet und wie eigentlich der Betrieb, also die Arbeit, das Handeln, nicht vertrauenswürdig. Wenn die Heimarbeit aber sogar auf dem privaten PC eines Mitarbeiters erfolgt, dann fällt auch das vertrauenswürdigem Gode weg.

REPORT | IT-SICHERHEITSSERVICES

Managed SOC: Angriffs-erkennung durch Dienstleister

Personal- und Ressourcenknappheit verhindern oft, dass Unternehmen in der Lage sind, die IT rund um die Uhr auf Anomalien und IT-Sicherheitsvorfälle zu überwachen. Dafür hat sich in den letzten Jahren ein ganzer Markt an Dienstleistern etabliert.

Von Patrick Bäuerle und Stefan Strobel



TRAC

- Wenn die eigene Organisation Opfer eines Hackerangriffs wird, muss man schnell reagieren. Die Zeit zwischen dem ersten Einbruch und der Veranschaulichung der kompletten Infrastruktur und der darauffolgenden Erpressung des Opfers ist in den letzten Jahren ständig kürzer geworden. Weniger als zwei Tage oder sogar nur wenige Stunden sind heute schnell reagierte. Die Zeit zwischen dem ersten Einbruch und der Veranschaulichung der kompletten Infrastruktur und der darauffolgenden Erpressung des Opfers ist in den letzten Jahren ständig kürzer geworden.
- Zahlreiche Dienstleister bieten Services zur Erkennung und Behandlung von Sicherheitsvorfällen an.
- Zwei wichtige Bausteine eines modernen Security Operations Center sind SIEM und EDR.
- Je nach deren Kombination lässt sich die Erkennung von Angriffen noch einmal verbessern – was sich in den Preismodellen der Anbieter niederschlägt.
- Um den passenden Anbieter für das eigene Unternehmen zu finden, ist zu definieren, welche Dienste nötig sind. Neben der genauen Analyse der eigenen Anforderungen ist auch eine kritische Auseinandersetzung mit den Angeboten der Anbieter notwendig.
- Nicht jedes Betriebsmodell passt zu jedem Unternehmen. KRITIS-Organisationen werden beispielsweise von Cloud-Diensten Abstand nehmen und On-Premises-Modelle vorziehen.

KI und Security: Schlagzeilen

Künstliche Intelligenz wird Zero-Day-Schwachstellen explodieren lassen

In today's modern workspace,
everyone has their own AI helper.

Indirect Prompt Injection Attacks: A Lurking Risk to AI Systems

Copilot's No-Code AI Agents Liable to Leak Company Data

Polizei warnt vor Betrugsmasche

Deepfake auf Facebook: Re Lockvogel für betrügerische

Betrüger haben ein Deepfake-Video von Reinhold Würth zu dubiosen Investments zu locken. Würth-Gruppe und der Masche.

Stand:
18.1.2026, 18:10 Uhr



Von Ulrike Schirmer

heise online > Künstliche Intelligenz > Anthropic launcht Claude Code Security – Cybersecurity-Aktien verlieren

Anthropic launcht Claude Code Security – Cybersecurity-Aktien verlieren

Das KI-Tool Claude Code Security von Anthropic analysiert Code kontextbasiert statt regelbasiert. Die Börse reagiert nervös, Aktienkurse geben nach.

🇬🇧 🔊 🖨️ 💬 12



(Bild: tadamichi/Shutterstock.com)

21.02.2026, 16:53 Uhr Lesezeit: 5 Min. | Developer

Von Matthias Parbel

Little Bay Beach, Anguilla (Getty Images)

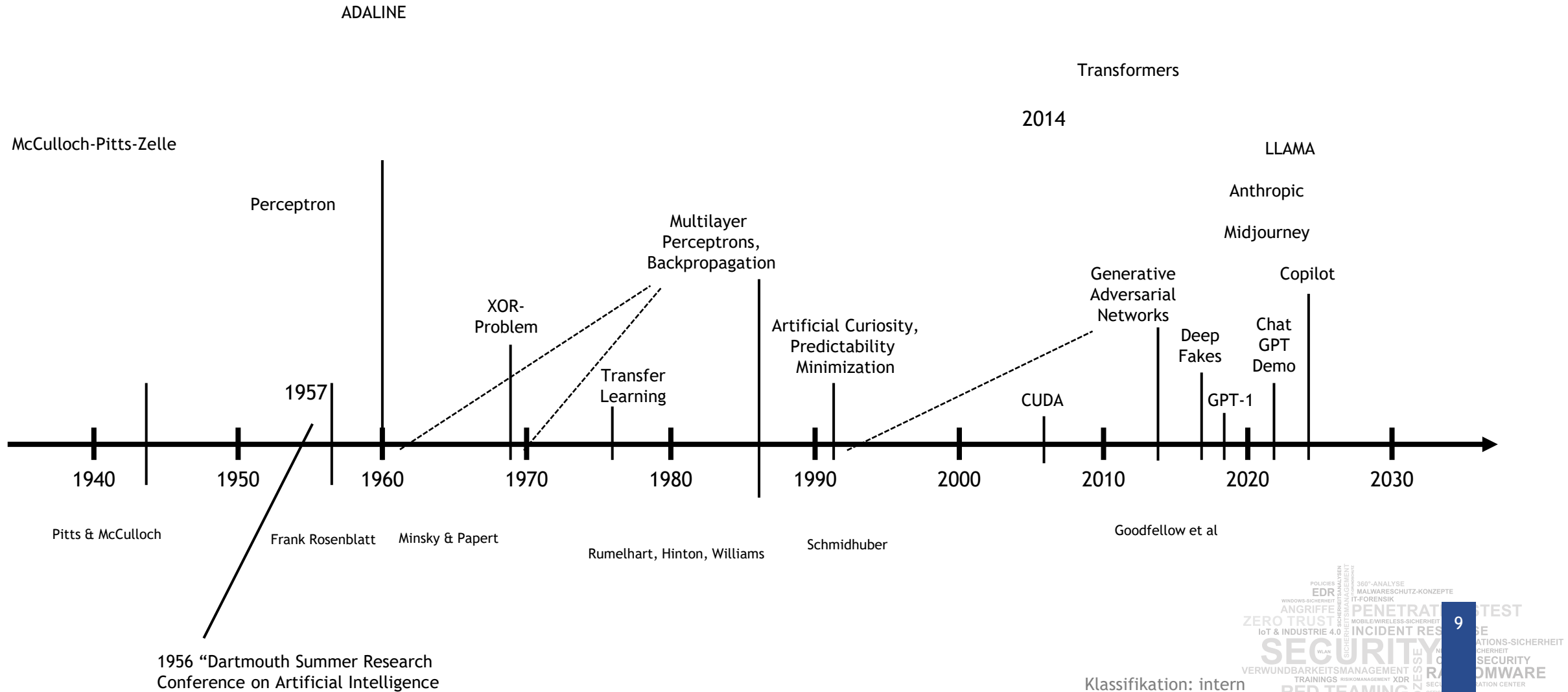
There are now more than 1 million “.ai” websites, contributing an estimated \$70 million to Anguilla’s government revenue last year

Data from Domain Name Stat reveals that the top-level domain originally assigned to the British Overseas Territory of Anguilla passed the milestone in early January.

Applying Agen
to Security Ope



Ausgewählte Meilensteine in der Geschichte der KI



Bereiche

Angriffe mit Hilfe von KI

Phishing &
Vishing

Schwachstellensuche

Malware-
Erstellung

Deep Fakes

Exploitentwicklung

Malware mit
eingebauter KI

Sicherheitsmaßnahmen mit KI

Bedrohungserkennung

Automatisierung

SOC-Unterstützung

Deep-Fake-Erkennung

Schwachstellen durch Nutzung von KI

Datenabfluss
(z.B. in Prompts)

Manipulation von KI-Applikationen
(z.B. Prompt Injection)

Angriffe auf interne Systeme und Daten
(z.B. Prompt Injection triggert Aktivitäten)

Schutzmaßnahmen gegen Schwachstellen bei KI-Nutzung

Absicherung eigener LLMs

Schutz vor Datenabfluss

POLICIES
EDR
WINDOWS-SICHERHEIT
ANGRIFFE
ZERO TRUST
IoT & INDUSTRIE 4.0
SICHERHEITSMANAGEMENT
IT-GRUNDSCUTZ
360°-ANALYSE
MALWARESCHUTZ-KONZEPTE
IT-FORENSIK
PENETRATIONSTEST
MOBILE/WIRELESS-SICHERHEIT
INCIDENT RESPONSE
APPLIKATIONS-SICHERHEIT
NETZWERKSICHERHEIT
CLOUD SECURITY
SECURITY
WLAN

KI als Werkzeug der Angreifer

Mit KI erzeugtes Bild

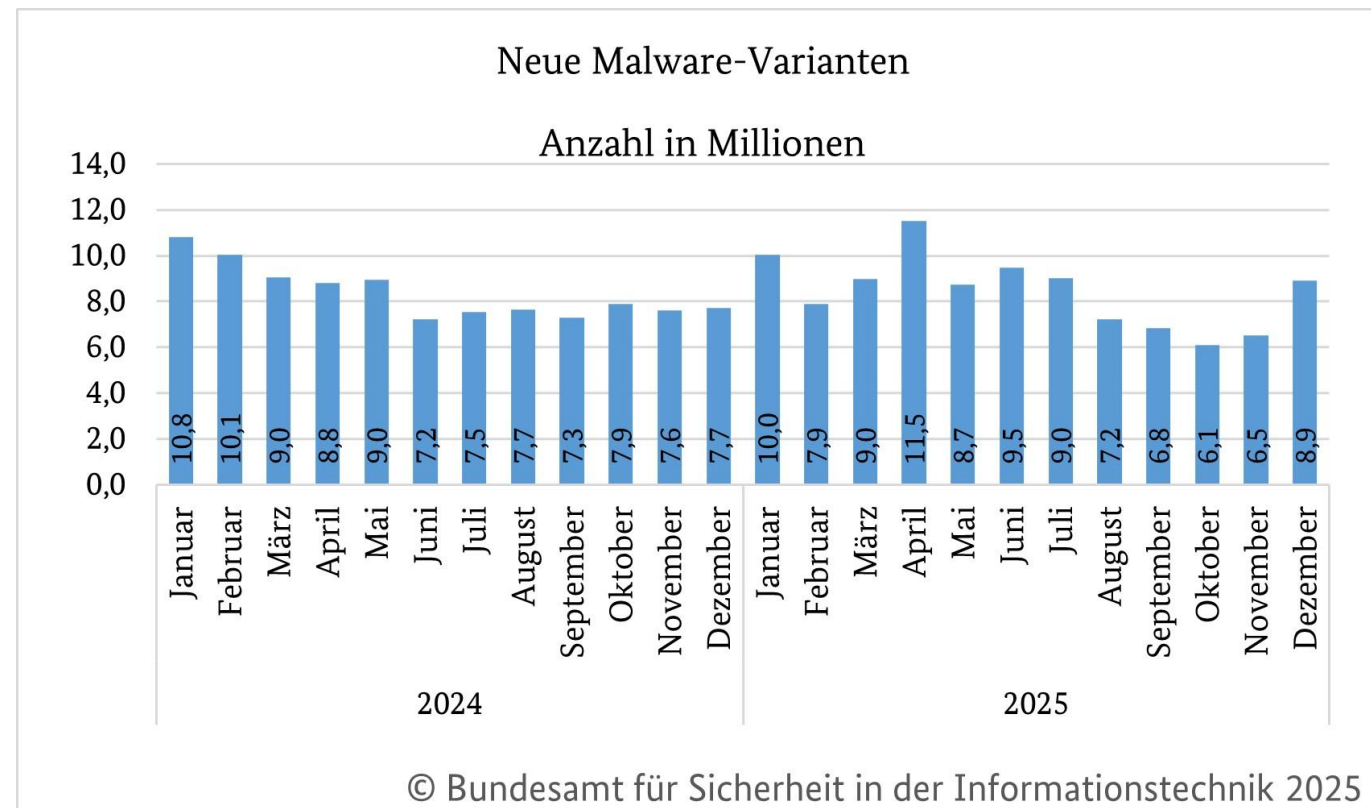


KI zur Entwicklung von Malware oder Exploits

- KI ist eine geschätzte Erleichterung für Software-Entwickler
 - Also auch für Entwickler von Malware
- Daraus entstehen jedoch keine genialen neuen Ansätze
 - Mehr Effizienz
 - Theoretisch mehr Variationen bereits bekannter Techniken
 - Das war aber auch bisher schon durch Malware-Builders, Obfuskierung etc. möglich
- Malware-Entwicklung wird einfacher, auch für inkompetente Angreifer
 - Durch MaaS war das schon ein Problem

Wo ist der Malware-Anstieg in den Statistiken?

- Zumindest bisher ist hier kein KI-Effekt zu sehen



Malware mit eingebauter KI „AI Powered“ Malware statt „AI Generated“

- **Phishing-Texte mit KI**
 - Sind nicht besser als Spear-Phishing-Texte, die sich ein Mensch ausdenkt
- **Aber**
 - können automatisch erzeugt werden
 - Können automatisch Kontext zum Opfer verarbeiten
 - Überwinden Sprach-Barrieren (Sprache des Opfers ist ggf. Fremdsprache für Angreifer)
 - Erstellung skaliert besser als manuelles Spear-Phishing
- **Beispiel Dynamit-Phishing von Emotet**
 - Damals automatisch erzeugte Antworten auf die Mails eines ersten Opfers
 - Oft mit Inkonsistenzen (z.B. „Du“ vs. „Sie“)
 - Mit KI erreicht man bessere Ergebnisse



Von Emotet zuvor auf infiziertem System ausgespähte authentische E-Mail

- Bertram Müller bekommt vermeintlich eine Antwort von Antje Meier
 - Mit einem Link auf Malware
- Der Computer von Antje Meier wurde infiziert und ihre Mails bzw. Kontakte wurden analysiert
 - Z. B. eine echte Frage von Bertram Müller



Bildquelle: CERT-bund

- [illegible]



Mit KI erzeugtes Bild

- Vishing ist Social Engineering über Sprache
 - typischerweise Telefon
- Der Angreifer ruft sein Opfer an und verleitet ihn zur gewünschten Aktion
 - Z. B. Preisgabe von Informationen, Anmeldung an einer Webseite, Ausführen einer Überweisung etc.
- Vorteil für den Angreifer:
 - Keine störenden E-Mail-Gateways
 - Fälschen einer Telefonnummer ist leichter als das Fälschen einer E-Mailadresse
 - Geringere Sensibilisierung der Mitarbeiter für die Gefahren durch Telefonanrufe
 - Direkte Interaktion



Klassifikation: intern

Künstliche Intelligenz Lösung für das „1-zu-1-Problem“



Künstliche Intelligenz
übernimmt die
Durchführung der Angriffe

Mit KI erzeugtes Bild

Klassifikation: vertraulich

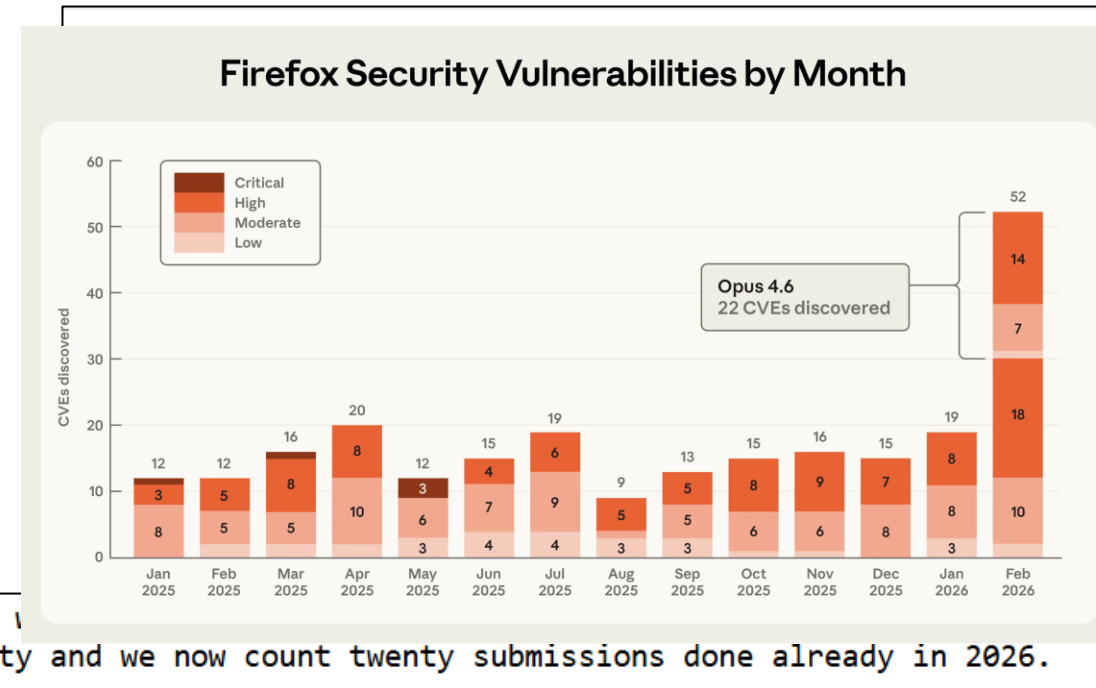
Wie kommt man an die passende Stimme?

- Mitschnitte von 10 Minuten reichen für eine hohe Qualität völlig aus
 - Aufgezeichnete Präsentationen
 - Vorgetäuschte Interviews
 - Mitschnitte von Teams-Calls unter Vorwand
 - Bezahlte Analysten-Calls

Suche nach Schwachstellen

- Fuzzing oder auch Quellcode-Analyse werden durch KI effizienter
 - Dadurch entstehen jedoch keine neuen Schwachstellen bzw. Angriffspunkte in Software
 - Die Programmierfehler müssen schon vorhanden sein
- Die Suche kann im Rahmen des SSDLC eingesetzt werden, aber auch von Angreifern
- Mit KI „gefundene“ unsinnige Schwachstellen als Problem für Hersteller und Bug Bounty
 - Beispiel curl und Log4J
- Suche scheint aber schnell besser zu werden
 - Siehe OpenSSL im Januar 2026, Mythos, ...

Klassifikation: intern



AISLE Discovered 12 out of 12 OpenSSL Vulnerabilities

Author: Stanislav Fort | Date Published: January 26, 2026

Subscribe to newsletter

Autonomous zero-day discovery in one of the most scrutinized codebases in the world

AISLE's autonomous analyzer found all 12 CVEs in the January 2026 coordinated release of OpenSSL, the open-source cryptographic library that underpins a substantial proportion of the world's secure communications. Some of these vulnerabilities had persisted in OpenSSL code for decades, evading the notice of thousands of security researchers.

Finding a genuine security flaw in OpenSSL is extraordinarily difficult. Even a single accepted vulnerability represents a rare achievement. The library's maturity and the community's vigilance make new discoveries exceptionally uncommon. This makes the January 2026 release an important milestone for autonomous security systems. As Tomáš Mráz, CTO of the OpenSSL Foundation, says,

January 2026: 20 reports

- February 2026: 13 reports

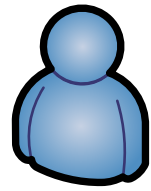
Neu ausgenutzte Schwachstellen in 2025 laut VulnCheck im Januar 2026



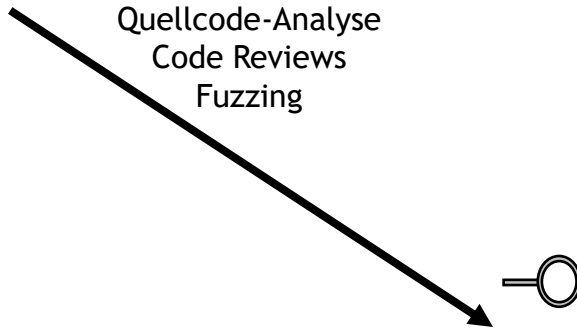
- 884 neue „Known Exploited Vulnerabilities“
 - 2025 zum ersten Mal ausgenutzt
- 28,73% bereits vor der Veröffentlichung ausgenutzt
- Mittlere Zeit von Veröffentlichung bis Ausnutzung bereits 2024 auf 5 Tage gesunken
 - Passt das zu SLAs mit Outsourcern?
- Netzwerk-Edge-Geräte am häufigsten angegriffen



KI zur Suche nach Software-Schwachstellen



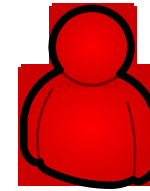
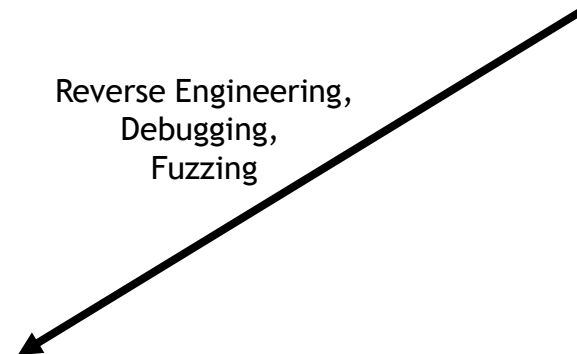
Quellcode-Analyse
Code Reviews
Fuzzing



Software



Reverse Engineering,
Debugging,
Fuzzing





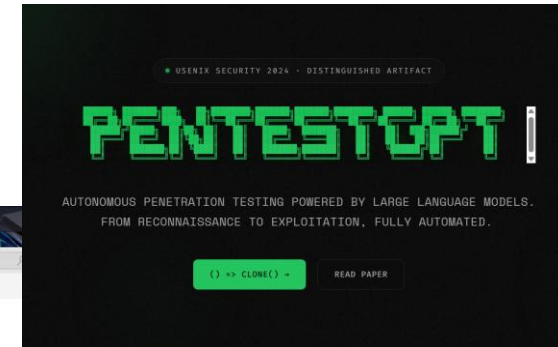
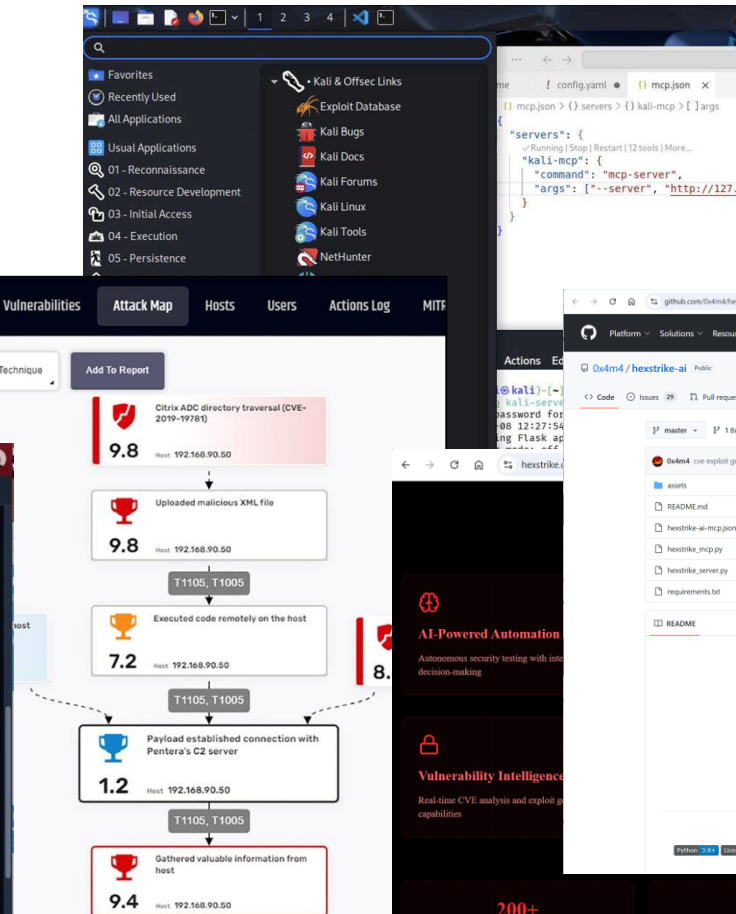
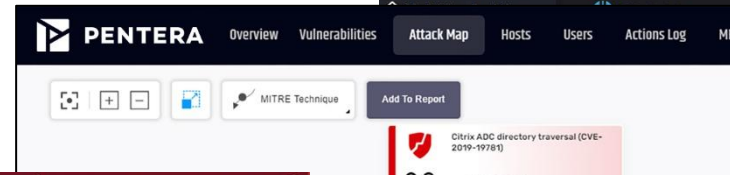
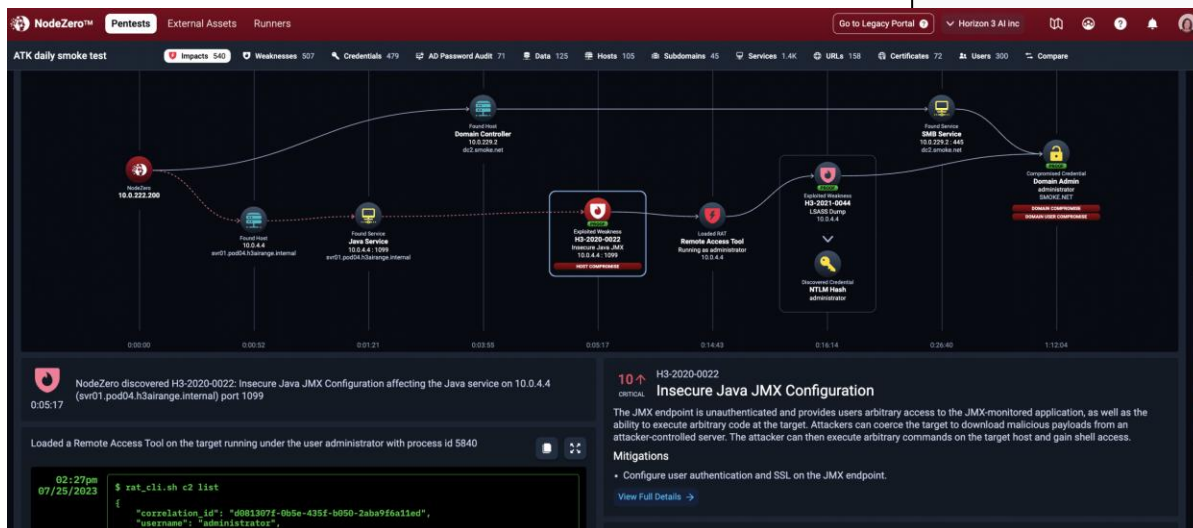
KI erzeugt Schwachstellen im Code

- Bei chinesischer KI wohl auch absichtlich ...

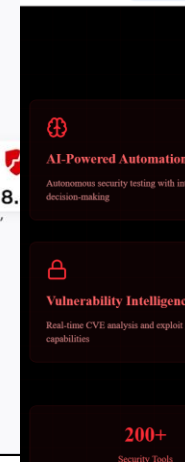
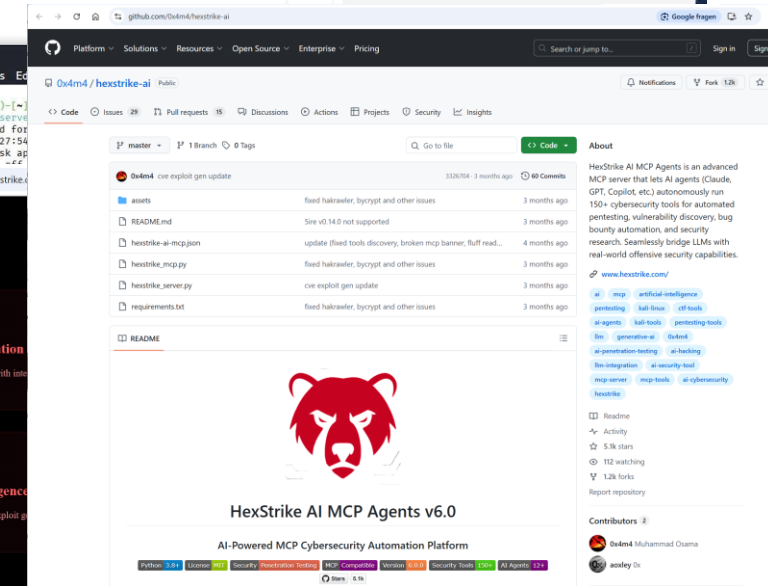
Technische Angriffe / Penetrationstests mit Hilfe von KI



- Automatisierte Suche und Ausnutzung von technischen Schwachstellen
- Open Source
 - Siehe iX 12/2025: Kali Linux mit GenAI und MCP (Jörg Riether), HexStrike
- Kommerzielle Produkte, z.B.
 - Pentera
 - Horizon3.ai



Das Ziel ist die Sammlung von etwaigen Logindaten oder eindeutigen möglichen Angriffsvektoren. Dazu habe ich einmal eine metasploitable VM etabliert, welche nachweislich dramatische Schwächen und Verwundbarkeiten hat. Führe an dieser Maschine bitte eine Analyse durch und schreibe mir einen sehr kurzen Report. Die IP lautet: 172.16.162.142. Fasse Dich bitte unbedingt sehr kurz und berichte mir ausschließlich von fundamentalen katastrophalen Funden. Bitte maximal 5-8 Sätze und ausschließlich in deutscher Sprache.



HexStrike AI MCP Agents v6.0

AI-Powered MCP Cybersecurity Automation Platform

200+ Security Tools

50+ AI Agents

100% Open Source

RED TEAMING 8 OFFICE

SE
ATIONS-SICHERHEIT
CHERNET
SECURITY
DMWARE
ATION CENTER

- Im Kontext Suche und Ausnutzung technischer Schwachstellen bringt KI keine neue Qualität, aber bessere Skalierbarkeit
 - Ein erfahrener Angreifer oder Pentester ist überlegen
 - Für Organisationen im Fokus der Angreifer ändert sich wenig
 - KI basierte Angriffe auf kleinere und schlechter gesicherte Organisationen werden jedoch einfacher und lohnend
 - Aber „von uns will doch keiner was“ hat auch bisher nicht funktioniert
 - Angriffe erfordern immer weniger technische Kenntnisse
 - Aber auch dieser Trend existiert schon länger, MaaS etc.
- KI kann Angriffe beschleunigen
 - Mit KI können Angriffe noch besser automatisiert werden als bisher
 - Die Zeit zwischen Veröffentlichung einer Schwachstelle, Initial Access und Erpressung wird weiter sinken



KI zur Verbesserung der Sicherheit

Mit KI erzeugtes Bild



KI / Machine Learning in der Security

- Sehr beliebt: Klassifikation bzw. Unterscheidung zwischen
 - Normalen Dateien und Malware
 - Typischer Anwendungsfall für „Supervised Learning“
 - Normalem und böartigen Verhalten auf Endgeräten
 - Normalen Mails und Spam
- Erkennen von Anomalien
 - Vergleich von Features mit als „normal“ gelernten Features
 - Typischer Anwendungsfall für „Unsupervised Learning“, z.B. für Netzwerk-Traffic
- uvm.

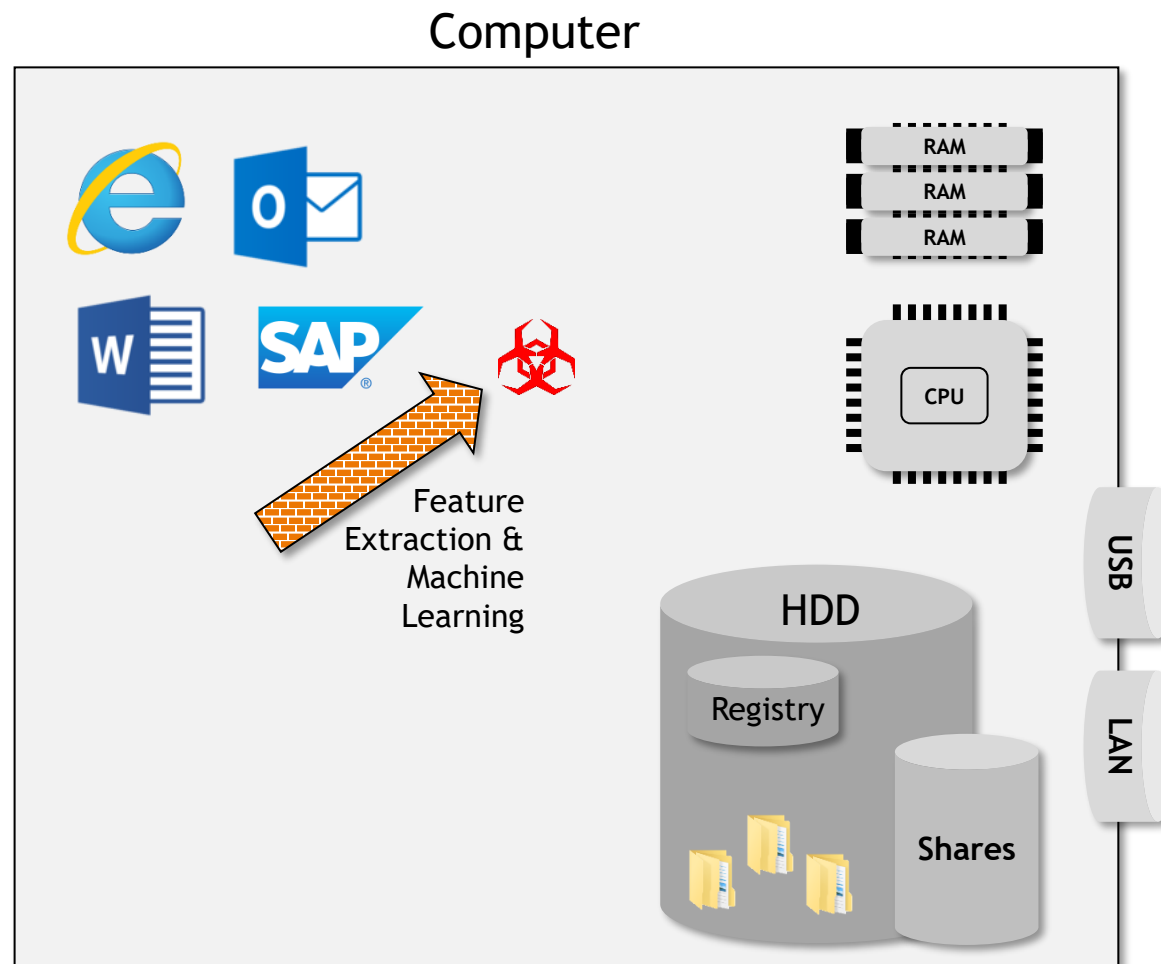
Hat Machine Learning das Spam-Problem gelöst?

Generelle Probleme mit KI in der Security

- Die zu klassifizierenden Objekte sind keine natürlichen Dinge
- Spam-Autor versucht aktiv die Erkennung zu verhindern
 - Text in Bildern
 - Falsche aber lesbare Schreibweise
 - Analyse des Modells und gezielte Umgehung
 - Massiv Spam-Mails als Non-Spam reporten -> Vergiften des Modells
 - usw.
- Machine Learning geht oft davon aus, dass man nicht versucht es zu umgehen
 - In der Security hat man oft einen intelligenten Gegner
- Weitere Herausforderung: Herkunft der Quelldaten für das Lernen
 - Gibt es genügend Quelldaten für alle Klassen?
 - Hat der Angreifer Einfluss darauf?

AV Reborn? Next-Generation AV mit KI?

- Feature-Extraktion mit Fokus auf exe-Dateien
- Neuronale Netzwerke beim Hersteller trainiert (Supervised)
 - Millionen von Malware-Samples
 - Millionen gutartige Dateien
- Ergebnis ist ein mathematisches Modell
 - Klassifikation von Dateien ist performant machbar
 - Feature-Extraktion, Berechnen des Modells
- Neue Malware oder Varianten gut erkennbar
 - Keine Signaturen nötig
- Updates nur noch alle 2-3 Monate



Bekannte Anbieter z.B.:

Cylance
Deep Instinct

Neuronale Netze vs. Signaturen

So einfach ist es nicht ...

- KI hat sicher Vorteile

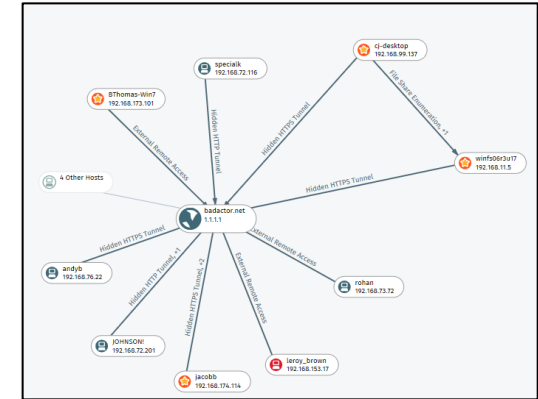
- Auch unbekannte Objekte können schnell klassifiziert werden
 - Bei signaturbasierten Verfahren muss der Hersteller erst eine neue Signatur erstellen und verteilen
 - Damit gehen meist Stunden oder sogar Tage verloren
 - Durch engere Cloud-Integration wird aber auch das schneller
- Weniger Performance-Verlust / Last auf den Endgeräten als bei langen Signaturlisten
 - Aber auch hier greift die Cloud (Reputation wird online abgefragt)

- Aber auch Nachteile

- Wenn das Objekt sich gegen die Erkennung wehrt
 - Bzw. der Autor die Erkennung verhindern möchte
 - Testen neuer Malware gegen eine Offline-Kopie der Sicherheitssoftware
 - Gezieltes Packen geht auch hier
- Bis dann ein neues neuronales Netz trainiert wurde, vergehen Wochen oder Monate

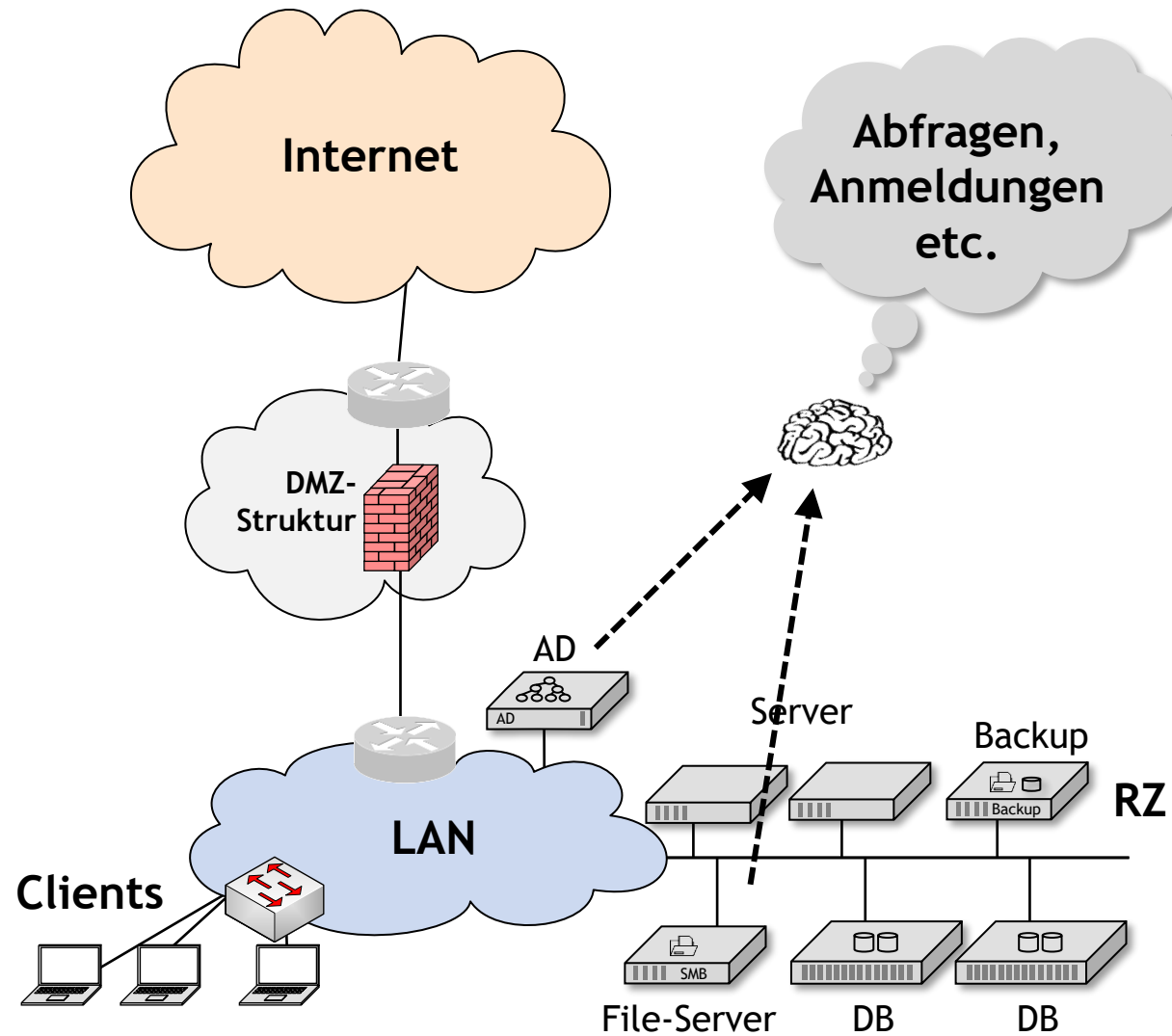
Fokus von KI basierter Malware-Erkennung

- Bei den meisten Produkten auf ausführbaren Programmen
 - PE-Files (primär exe)
- Diese werden aber immer häufiger schon auf anderem Weg blockiert
 - Mail-Programme / Gateways blockieren ausführbare Attachments
 - Whitelisting (z.B. AppLocker) blockiert Programme aus Verzeichnissen, in die der Anwender schreiben kann
- Der Mehrwert solcher Lösungen schwindet daher, sofern nicht andere Dateitypen ebenfalls mit KI unterstützt werden

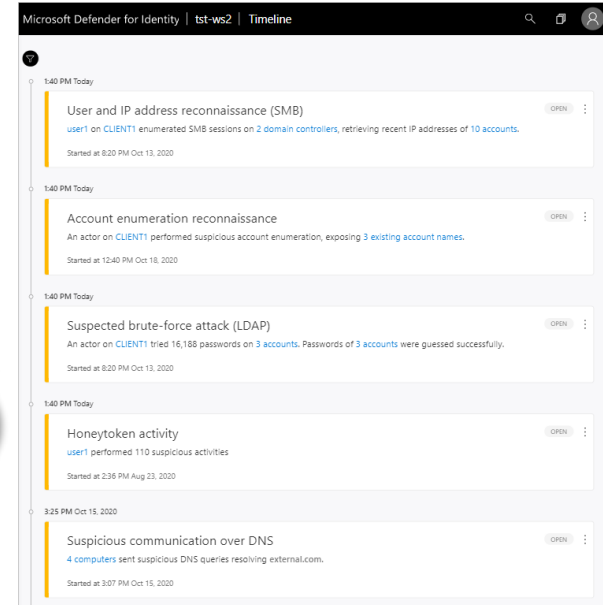


Klassifikation: intern

Alternative Ansätze zur Erkennung: AD-Security



Bekannte Anbieter z. B.:
Microsoft, Netwrix,
Tenable, CyberArk



SIEM auf Basis von KI



← → ↺ ↻

sentinelone.com/de/platform/ai-siem/

☆

Ein Leader im Gartner® Magic Quadrant™ für Endpoint Protection Platforms 2025. Seit fünf Jahren in Folge. [Bericht lesen →](#)

Sie haben eine Sicherheitsverletzung festgestellt? [Blog](#) [🔍](#) [🌐](#) [▼](#)

Plattform [▼](#) Warum SentinelOne? [▼](#) Services [▼](#) Partner [▼](#) Ressourcen [▼](#) Unternehmen [▼](#)

[Demo anfordern](#) [Kontakt](#)

Singularity AI SIEM für autonome SOC's

Schützen Sie Ihr gesamtes Unternehmen mit der branchenweit schnellsten KI-gestützten offenen Plattform für alle Ihre Daten und Workflows, die auf SentinelOne Singularity™ Data Lake aufbaut.

Plattform Products Solutions Resources Partners Company

[Request a Demo →](#)

AI SOC ANALYST

ELEVATE ANALYSTS WITH A SMARTER, AI-POWERED SOC

Transform your SOC with a fully autonomous AI analyst. Supercharge your team's speed, efficiency, and resilience through instant, AI-powered triage, investigation and response—reducing MTTR by 83%.
Unleash the power of always-on defense that learns from human expertise, never tires, and consistently delivers evidence-driven results so your analysts can focus on evolving complex threats, not repetitive triage.

[Watch a Demo →](#) [Request a Demo →](#)

Gartner 2025 Magic Quadrant for SIEM—Gurucul Named a Leader!

Get the latest Gartner report and learn how Gurucul is ushering in the new era of Next-Gen SIEM, an evolution required for the smarter SOC.

[Get the Report in Your Inbox!](#)

Lumi AI Guru

Hi there! I'm Lumi, an AI Guru from Gurucul. I noticed you checked out our AI SOC Analyst product. Any questions about how it can fit into your current needs?
Saw that you're interested in our AI SOC Analyst! I'm available if you have questions or want to chat.

[Schedule a meeting](#)

[Ask a question](#)

AUTOMUNBIAS ALERT

The AI SOC Analyst acts as a virtual L1 analyst, working 24/7 to detect and respond to threats.

← → ↺ ↻

exabeam.com/explainers/siem/what-is-siem

☆

Exabeam Intro

Table of Contents

State of the SIEM

5 Key Benefits

14 Core SIEM Features

Comparing SIEMs

Cybersecurity SIEM Use Cases

Exabeam Future

← → ↺ ↻

crowdstrike.com/en-us/platform/next-gen-siem/

☆

AI Summit: Accelerating Secure AI Adoption and Development [Register now](#)

Experienced a breach? | [Blog](#) | [Contact us](#) | [Careers](#) | [Latest Innovations](#)

Plattform [▼](#) Services [▼](#) Solutions [▼](#) Why CrowdStrike [▼](#) Resources [▼](#) Pricing [▼](#)

[Start free trial →](#)

Next-Gen SIEM [Explore ▼](#)

Platform / Next-Gen SIEM

SOCByte

ALL BYTES SECURED

Platform [▼](#) For MSSPs [▼](#) Why SOCByte? [▼](#) Partners [▼](#) Blogs [▼](#)

[Request a demo →](#)

SOCByte SIEM

SOCByte SIEM helps your team centralize logs, correlate events, and respond in real time, powered by AI. Fully customizable, parser-friendly, and built for hybrid SOC's.

[Request a Demo](#)

Investigate with Charlotte AI

Summary

Charlotte AI

Between 2024-04-02 at 17:00:00 UTC and 2024-04-02 at 17:00:00 UTC, 10 events were identified on host 10.10.10.10. The events were identified on host 10.10.10.10, indicating potential use of Command and Control and Security Information Systems.

On 2024-04-02 at 17:00:00 UTC, a system breach event occurred on the host 10.10.10.10 and 10.10.10.10. The events were identified on host 10.10.10.10, indicating potential use of Command and Control and Security Information Systems.

Investigating this event, a series of suspicious activities were detected on the same host. Multiple instances of PowerShell were executed, some with highly elevated and dangerous commands, suggesting potential malicious activity. These events align with known MITRE ATT&CK tactics and techniques, including: Execution, Persistence, and Command and Control. Suspicious network connections were observed, including connections to external IP addresses. These connections were identified through network logs and network analysis. Additionally, there were indications of a possible malware infection. A PowerShell script was executed, which appears to be a malicious script. The execution of this script and the subsequent network connections suggest a possible compromise of the system. The events were identified on host 10.10.10.10, indicating potential use of Command and Control and Security Information Systems.

[Add to notes](#)

[Report](#)

AI Agents im SOC

RADIANT

Platform Solutions Customers Company Resources

Scale and optimize your SOC with Agentic AI



Conifers

Product Enterprise MSSPs / MDRs Resources Company

Your SOC AI Force Multiplier

Conifers is the AI SOC agents platform that delivers deep, contextual investigations, adapted to your own data, decisions, and risk tolerance.

By continuously learning and adapting, and with the smart use of AI, Conifers scales your SOC effectiveness and efficiency—becoming a true force multiplier for your team.

Your assets

prophet

Product Customers Resources Why Prophet AI Company Blog

Request a Demo

A Smarter, Scalable SOC Powered by Agentic AI

Transform how you detect, investigate, and respond with AI as a SOC force multiplier

Request demo

Trusted by leading innovators:

IMPASS Moveworks SPOTNANA Clari Zip instacart

Dropzone AI

Product Solutions Customers Pricing Partners Careers Company Resources

AI SOC Analysts that never sleep so you can eliminate alert backlogs

The Dropzone AI SOC analyst replicates the techniques of elite analysts to autonomously investigate and solve every alert. Deploys in minutes.

Self-Guided Demo Request a Demo



Chosen by Forward-Thinking Security Teams

zapier Fortune 500 IT Services Provider UiPath Mysten Labs ECS

torq

Platform Solutions Company Resources Customers

Login Get a Demo

The World's First Multi-Agent System for the SOC

AI Agents for the SOC

Scan observables for this case



Achieve unprecedented SecOps efficiency and precision with AI Agents for the SOC that collaborate, plan, and reason to autonomously analyze and resolve security threats.

Oevlar AI

Product Solutions Resources Company

Request a demo

Automated Alert Investigation. No Playbooks.

Rule-based detection is no match for today's threats. Level-up with AI SOC Analysts.

Book a demo

3 min average time to investigate alerts

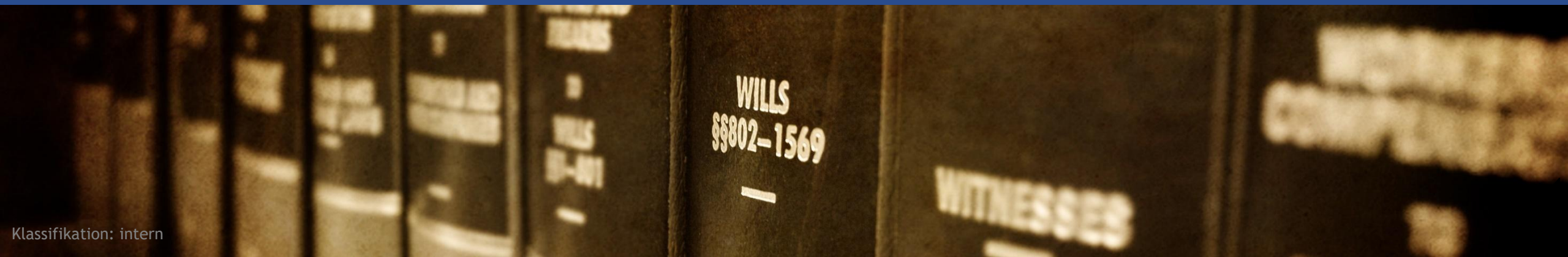
Up to 80% of tickets closed automatically

24/7 nonstop investigations

100% happier SOC analysts



Neue Probleme und Schwachstellen durch bzw. in KI



Angriffe auf KI-Erkennung (2015)



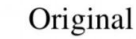
x

“panda”

57.7% confidence

- Gezielte Angriffe auf KI-Erkennung / neuronale Netze
- Auch auf Malware-Erkennung übertragbar

Original



Custom Sign Attacks

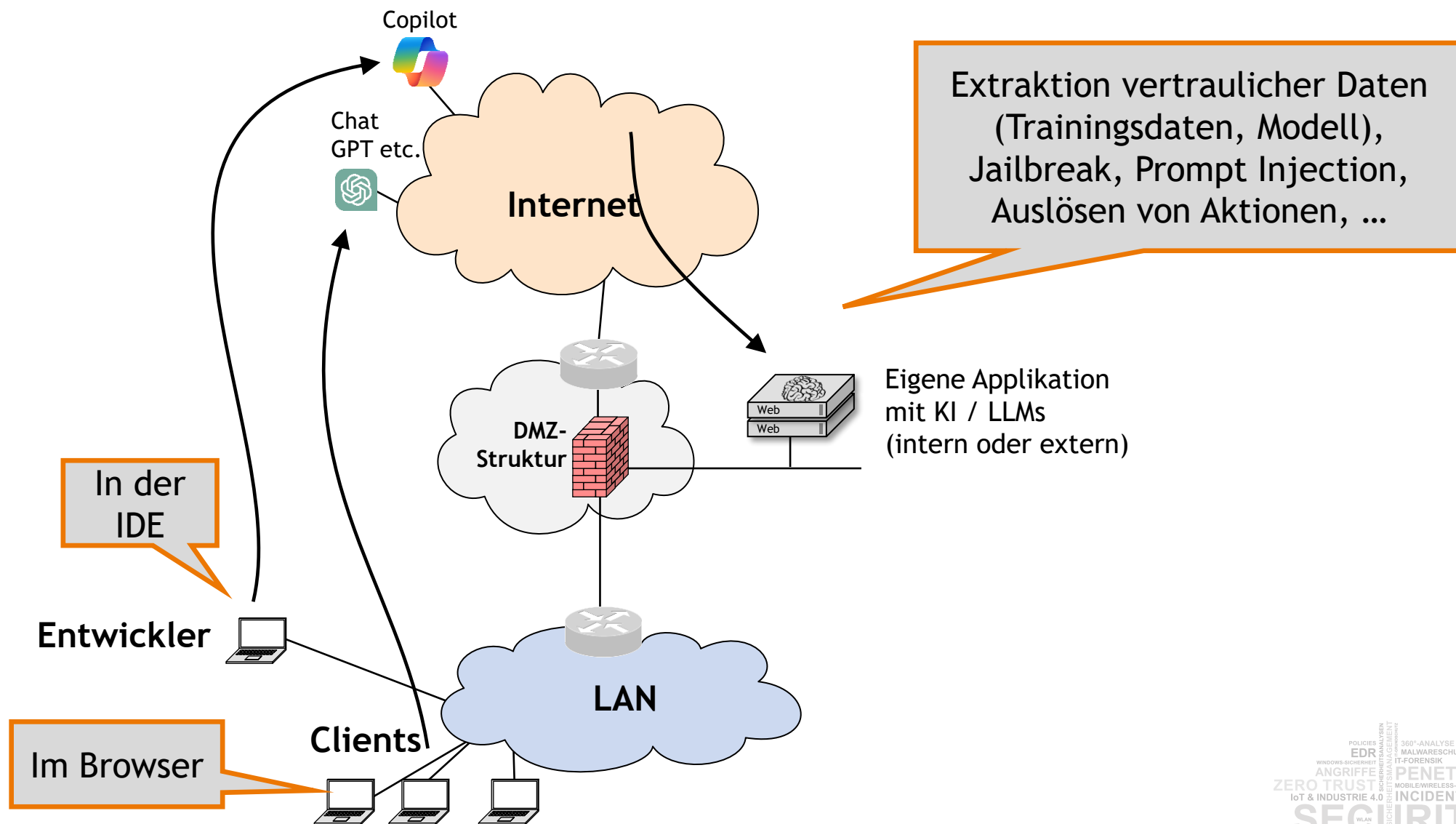


DARTS: Deceiving Autonomous Cars with Toxic Signs, Princeton University / Purdue University



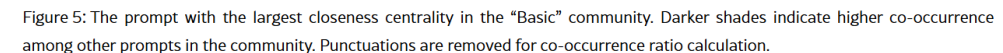
-

Angriffe auf eigene KI







¹CISPA Helmholtz Center for Information Security ²NetApp



Weiterentwicklung erhöht die Risiken

Früher:
KI reagiert auf Fragen und liefert Text, Grafiken etc.

 ChatGPT ▼ Anmelden Kostenlos registrieren ?

erstelle eine kurze Zusammenfassung der Risiken durch die Nutzung von KI

Die Nutzung von Künstlicher Intelligenz (KI) birgt mehrere Risiken, darunter:

- 1. Arbeitsmarktveränderungen:** KI kann viele Arbeitsplätze automatisieren, was zu Arbeitsplatzverlusten und sozialen Ungleichheiten führen könnte.
- 2. Verzerrung und Diskriminierung:** KI-Systeme können bestehende Vorurteile und Diskriminierungen in den Daten, mit denen sie trainiert wurden, übernehmen, was zu unfairen Entscheidungen führt (z.B. bei der Kreditvergabe oder Personalrekrutierung).
- 3. Datenschutz und Sicherheit:** KI erfordert große Mengen an Daten, was Bedenken hinsichtlich des Datenschutzes und der Sicherheit aufwirft. Missbrauch oder unzureichender Schutz von sensiblen Daten sind Risiken.

Agentic AI erstellt einen Plan und führt Aktionen aus

Streit zwischen Anthropic und der US-Regierung laut Reuters im Februar



World ▾ Business ▾ Markets ▾ Sustainability ▾ More ▾

Exclusive: Pentagon clashes with Anthropic over military AI use, sources say

By Deepa Seetharaman, David Jeans and Jeffrey Dastin

January 29, 2026 10:19 PM GMT+1 · Updated January 30, 2026

WEAPONS TARGETING

In its discussions with government officials, Anthropic representatives raised concerns that its tools could

AI

Research Economic Futures

Announcements

Statement from Dario Amodei on our work with the Department of War

26. Feb. 2026

I believe deeply in the existential importance of using AI to defend the United States and other democracies, and to defeat our autocratic adversaries.

heise online > Künstliche Intelligenz > Anthropic > Anthropic-CEO nennt Vorgehen des Pentagons „vergeltend und strafend“

Anthropic-CEO Amodei wehrt sich in einem Interview gegen die Einstufung als Sicherheitsrisiko und beruft sich auf amerikanische Grundwerte.



Anthropic-Gründer Dario Amodei (Bild: Anthropic)

TECHNOLOGY

Altman says OpenAI agrees with Anthropic's red lines in Pentagon dispute

BY JULIA SHAPERO - 02/27/26 12:02 PM ET

FOLLOW ON Google News



human oversight, some of the

ry 9 department memo on AI

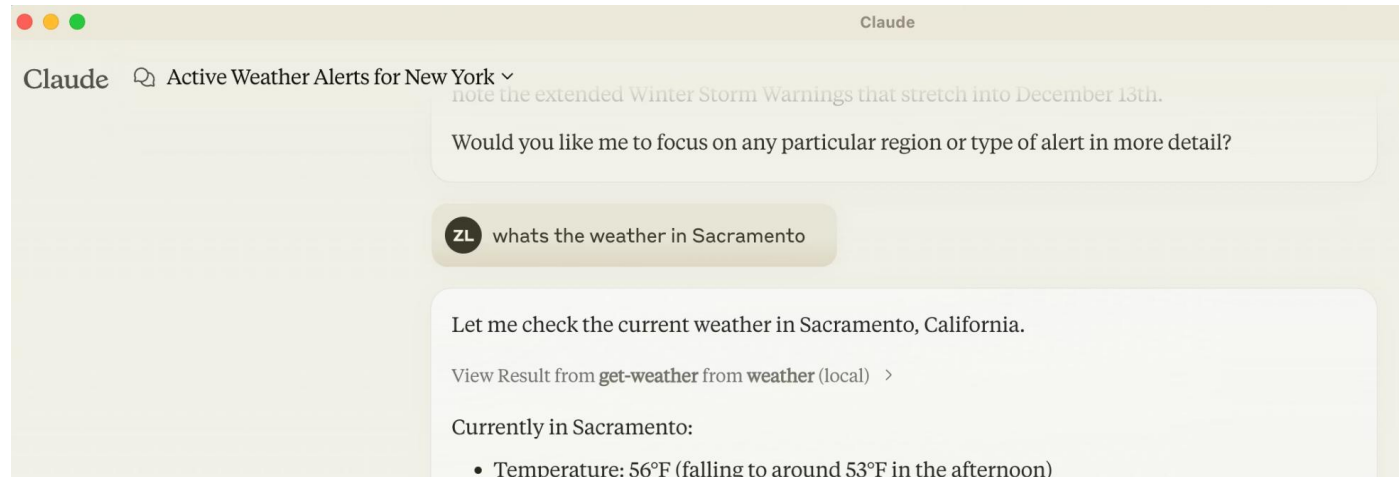
mercial AI technology

w, sources said.

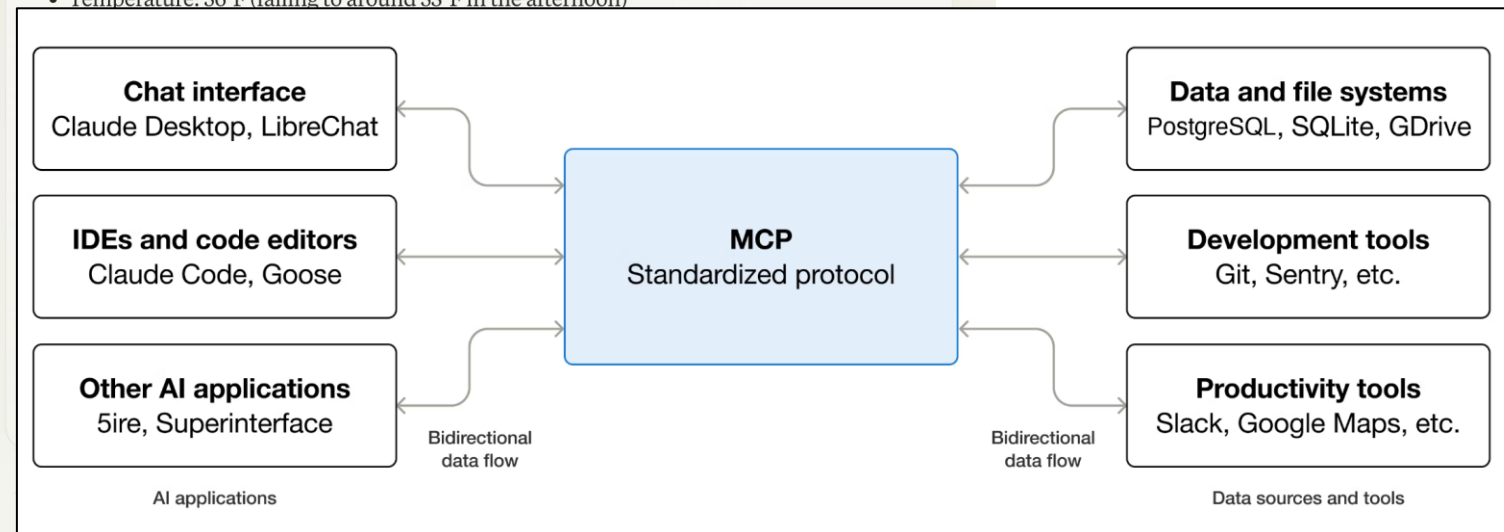


SECURITY

Model Context Protocol - MCP



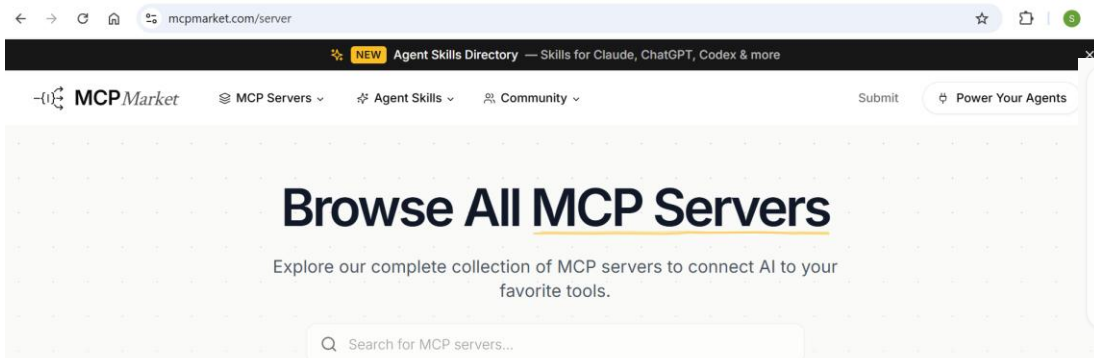
- 2024 von Anthropic vorgestellt
- „Agentische“ KI
- Erlaubt der KI zu interagieren
 - Daten abfragen
 - Aktionen steuern



Bildquelle: <https://modelcontextprotocol.io/docs/develop/build-server>

Klassifikation: intern

MCP für alles



Chrome Browser

★ 9.844

Exposes Chrome browser functionality to AI assistants for complex browser automation,...



Excel

★ 3.095

Enables Excel file manipulation capabilities without requiring Microsoft Excel installation.



Google Workspace

★ 1.080

Integrates Google Workspace services with AI assistants and applications via the Model...



Ghidra

★ 7.012

Enables Large Language Models (LLMs) to autonomously reverse engineer application...



Slack

★ 1.079

Connects Slack workspaces to Model Context Protocol (MCP) for enhanced...



MySQL

★ 1.040

Enables LLMs to inspect MySQL database schemas and execute read-only queries.



GitHub

Enables advanced automation and interaction capabilities with GitHub APIs for developers...

API Development



Task Master

★ 24.690

Streamline AI-driven development workflows by automating task management with Claude.

Developer Tools



FastMCI

Facilitates the cr Protocol (MCP)...

Other

wright

★ 25.239

II ElevenLabs

★ 1.139

Enables interaction with Text to Speech and audio processing APIs for MCP clients.



GPT Researcher

★ 24.774

Database Management



HexStrike AI

★ 5.769

Empowers AI agents with autonomous offensive cybersecurity capabilities by...



Kubernetes

★ 967

Provides a Model Context Protocol (MCP) server for managing Kubernetes and...



MCP Schwachstellen

- MCP erhöht nicht nur den Impact von Angriffen auf KI
 - Aktionen statt nur Textausgabe
- MCP bringt zahlreiche neue Schwachstellen

CVE-2025-6514 Detail

AWAITING ANALYSIS

QUICK IN

Description

mcp-remote is exposed to OS command injection when connecting to the authorization_endpoint response URL

Metrics

CVSS Version 4.0 CVSS Version 3.x CVSS Version 2.0

NVD enrichment efforts reference publicly available information to associate vector strings. CVSS information contributed by other sources is also displayed.

CVSS 3.x Severity and Vector Strings:



NIST: NVD

Base Score: N/A

NVD assessment not yet provided.



CNA: JFrog

Base Score:

9.6 CRITICAL

Vector:

CVSS:3.1/AV:N/AC:L/PR:N/UI:R/S:C/C:H/I:H/A:H

Klassifikation: intern

New Docker Hardened Images are now FREE for every developer. Register f

RESEARCH REPORT

State of AI Agent Security 2026

We audited how 30 popular AI agent projects handle authorization. The results are alarming: nearly all of them rely on unscoped API keys with no per-agent identity, no user consent, and no revocation mechanism.

Published March 15, 2026 · By Sanjeev Mishra · Grantex Research

93%

Use unscoped API keys

0%

Have per-agent identity

97%

No user consent flow

100%

No per-agent revocation

We Scanned 1,808 MCP Servers. 66% Had Security Findings.

March 14, 2026



AgentSeal Team

14 min read · 1552

11

6. Secret Exposure and Credential Theft



Produkte zur Absicherung der neuen von KI verursachten Probleme







Fazit



Wie bei jeder neuen IT-Technik

- Kann die neue Technik verwendet werden, um die Sicherheit zu verbessern
- Kann sie als Werkzeug von Angreifern verwendet werden
- Bringt die Technik selbst neue Schwachstellen mit sich
- So war / ist es auch mit Web, mit mobilen Geräten, mit der Cloud etc.



Mit KI erzeugtes Bild