



Holger Reibold

Autonomous Red Teaming mit KI

Von manuellen Pentests zu
autonomen Angreifen:
KI-basierte Systeme für
kontinuierliche und adaptive
Sicherheitsüberprüfung

BRAIN-MEDIA.DE

Policy & Governance Template

Der Einsatz autonomer Red-Teaming-Systeme erfordert klare, verbindliche Rahmenbedingungen, um Sicherheit, Kontrolle und Nachvollziehbarkeit zu gewährleisten. Während die technologische Leistungsfähigkeit dieser Systeme erheblich ist, entsteht ihr tatsächlicher Mehrwert erst durch eine strukturierte Einbettung in organisatorische und regulatorische Kontexte.

Das vorliegende Template dient als praxisorientierte Grundlage zur Definition von Policies und Governance-Strukturen. Es unterstützt Organisationen dabei, den Einsatz autonomer Angriffssimulationen klar zu begrenzen, Verantwortlichkeiten festzulegen und technische sowie operative Kontrollmechanismen zu etablieren. Ziel ist es, ein Gleichgewicht zwischen effektiver, realitätsnaher Sicherheitsüberprüfung und kontrolliertem, risikoarmem Betrieb zu schaffen.

Das Template ist bewusst modular aufgebaut und kann an unterschiedliche Anforderungen, Branchen und Reifegrade angepasst werden. Es bildet damit einen Ausgangspunkt für die Entwicklung unternehmensspezifischer Richtlinien und stellt sicher, dass autonome Systeme nicht nur leistungsfähig, sondern auch steuerbar und regelkonform eingesetzt werden.

1. Zielsetzung

Diese Policy definiert Rahmenbedingungen für den Einsatz von Autonomous Red Teaming zur kontinuierlichen Sicherheitsüberprüfung. Ziel ist die realitätsnahe Simulation von Angriffsszenarien unter Einhaltung technischer, organisatorischer und regulatorischer Vorgaben.

2. Geltungsbereich (Scope)

In Scope:

- Interne IT-Systeme
- Cloud- und Hybrid-Infrastrukturen
- Identitäts- und Zugriffsmanagement-Systeme

Out of Scope:

- Kritische Kernsysteme ohne explizite Freigabe
- Systeme Dritter ohne vertragliche Zustimmung

3. Zulässige Aktivitäten

- Automatisierte Reconnaissance und Kontextanalyse
- Schwachstellenerkennung (passiv und aktiv, nicht disruptiv)
- Kontrollierte Exploitation (nur innerhalb definierter Grenzen)
- Living-off-the-Land-Techniken mit Systembelastung

Nicht zulässig:

- Destruktive Aktionen
- Denial-of-Service-ähnliche Aktivitäten
- Privilegieneskalaation außerhalb definierter Grenzen

4. Rollen und Verantwortlichkeiten

- Security Lead: Gesamtverantwortung, Freigaben etc.
- Red Team Operator: Steuerung und Überwachung
- System Owner: Freigabe von Scopes

- Governance/Compliance: regulatorischer Anforderungen

5. Kontrollmechanismen

- Policy-basierte Entscheidungslogik für alle Agenten
- Definierte Kontrollpunkte bei kritischen Aktionen
- Echtzeit-Monitoring aller Aktivitäten
- Automatisierte Limitierung von Aktionen

6. Kill Switch & Intervention

- Hard Stop: Sofortige Beendigung aller Aktivitäten bei kritischen Ereignissen
- Soft Intervention: Anpassung oder Pausierung einzelner Aktionen
- Auslöser: Unerwartete Systeminstabilität, Überschreitung definierter Schwellenwerte, manuelle Intervention durch autorisierte Rollen

7. Logging & Audit

- Vollständige Protokollierung aller Aktionen und Entscheidungen
- Speicherung in revisionssicherem System
- Nachvollziehbarkeit von Kontext, Ziel und Ergebnis jeder Aktion

8. Datenschutz & Compliance

- Minimierung der Datenerhebung auf notwendige Informationen
- Schutz sensibler Daten durch Zugriffsbeschränkungen
- Einhaltung geltender regulatorischer Anforderungen

9. Review & Anpassung

- Regelmäßige Überprüfung der Policy (mind. quartalsweise)
- Anpassung bei Änderungen der Infrastruktur oder Bedrohungslage
- Dokumentation aller Änderungen

Dieses Template dient als Ausgangspunkt und muss an die spezifischen Anforderungen, Risiken und regulatorischen Rahmenbedingungen der jeweiligen Organisation angepasst werden.

Mehr zum Thema Autonomous Red Teaming mit KI

Der vollständige Leitfaden „Autonomous Red Teaming mit KI“

 [Jetzt bei Amazon bestellen](#)