

Risikobasiertes Framework für robuste LLM-Integration

Sandro Hartenstein, M.Sc. sandro.hartenstein@hwr-berlin.de



MOTIVATION

Semantische Fragilität, Inkonsistenzen und Halluzinationen führen häufig zu Fehlentscheidungen in LLM-basierten Systemen. Trotz vorhandener Leitlinien fehlt eine praxisnahe Operationalisierung zur systematischen Absicherung derartiger Integrationen.

METHODIK

Design Science Research

PROBLEM
Gap-Analyse

SOLUTION
Framework

DESIGN
Threat Modeling, Controls

EVALUATION
Fallstudien, Tests

COMMUNICATION
Publikation, Tools

Empirische Erhebungen:
4289 OpenAPI Specs • 430 Adversarial Tests • 31 Online-Umfrage • 10 Experten-Diskussion

ERKENNTNISSE — Praxisprobleme & Robustheitslücken

Top-Praxisprobleme
(Online-Umfrage, N=31)

Semantische Fragilität

3.24/5

Halluzinationen

3.14/5

Komplexe Anweisungen

2.90/5

Inkonsistenz

2.83/5

Folgen:

- 41.9% Funktionale Fehler
- 41.9% Manuelle Korrekturen
- 29% Fehlentscheidungen

Provider-Robustheit
(Adversarial Testing)

Amazon Titan

87.6%

Azure GPT-4

82.1%

IBM Watsonx

82.1%

Google Gemini

66%

Kernerkenntniss:

• 21.6% Provider-Varianz (66–87.6%)

Rahmendaten:

- 430 Test-Prompts (promptfoo)
- Zeitpunkt: Juli 2024

SOLUTION & DESIGN

Risk Assessment Matrix

Monitoring Toolkit

Test Framework

Mitigation Catalog

Robustness Checklist

Integrationsleitfäden

EVALUATION

Fallstudien

Experten-Review

Adversarial Tests

GQM-Metriken

ZIEL

Ein risikobasiertes Framework für robuste LLM-Integration: Empirische Erkenntnisse werden mit bestehenden Standards verknüpft, um durch Threat Modeling, priorisierte Controls und praxisnahe Leitfäden die Verlässlichkeit von LLM-basierten Systemen zu erhöhen.

Details & Quellen

<https://ha.rtenstein.com/esapi.html>