



Holger Reibold

Lokale KI unter Kontrolle

Sichere Architektur,
Betrieb und Governance von
GenAI- und RAG-Systemen

BRAIN-MEDIA.DE

Holger Reibold

Lokale KI unter Kontrolle

Sichere Architektur, Betrieb und
Governance von GenAI-
und RAG-Systemen

BRAIN-MEDIA.DE

Alle Rechte vorbehalten. Ohne ausdrückliche, schriftliche Genehmigung des Verlags ist es nicht gestattet, das Buch oder Teile daraus in irgendeiner Form durch Fotokopien oder ein anderes Verfahren zu vervielfältigen oder zu verbreiten. Dasselbe gilt auch für das Recht der öffentlichen Wiedergabe. Der Verlag macht darauf aufmerksam, dass die genannten Firmen- und Markennamen sowie Produktbezeichnungen in der Regel marken-, patent- oder warenrechtlichem Schutz unterliegen.

Verlag und Autor übernehmen keine Gewähr für die Funktionsfähigkeit beschriebener Verfahren und Standards.

© 2026 Brain-Media.de

ISBN: 978-3-95444-314-7

Cover: Freepik

Brain-Media.de

Dr. Holger Reibold – Hubert-Müller-Str. 52c – 66113 Saarbrücken

info@brain-media.de – www.brain-media.de

Inhaltsverzeichnis

- Inhaltsverzeichnis I
- Vorwort 1
- 1 Private KI im Infrastrukturkontext 5
 - 1.1 Vom Experiment zur produktiven Plattform 7
 - 1.2 Treiber für On-Prem-GenAI 10
 - 1.3 Typische Einsatzszenarien für lokale KI 13
 - 1.4 Architekturentscheidungen 15
 - 1.5 Ökonomisches und Organisatorisches 19
 - 1.6 Management Summary 22
- 2 Bedrohungsmodell für lokale GenAI 23
 - 2.1 Neue Angriffsflächen durch LLM-Systeme 25
 - 2.2 Veränderte Anwendungssicherheit 29
 - 2.3 Prompt Injection und Jailbreaks 31
 - 2.4 Datenvergiftung und Retrieval-Angriffe 33
 - 2.5 Risiken beim Fine-Tuning 35
 - 2.6 Model Extraction und Abuse-Muster 36
 - 2.7 Ableitung eines praxisnahen Threat Models 38
 - 2.8 Management Summary 41

3	Architekturprinzipien.....	43
3.1	Sicherheitsziele und Designprinzipien.....	44
3.2	Trust Boundaries systematisch festlegen.....	47
3.3	Sicherheitszonen und Default-Deny.....	49
3.4	Entkopplung von Orchestrierung	51
3.5	Identitätsbasierte Dienstkommunikation.....	55
3.6	Referenzmuster	57
3.7	Management Summary	59
4	Die sichere RAG-Referenzarchitektur	61
4.1	Gesamtüberblick und Komponenten	64
4.2	Gateway und Zugriffsschicht.....	68
4.3	Orchestrierung und Prompt-Verarbeitung	70
4.4	RAG-Pipeline und Vektorstore.....	72
4.5	Model Serving und Embedding-Dienste	75
4.6	Kritische Datenflüsse und Architekturfehler.....	77
4.7	Air-Gapped-Deployment-Strategien	79
4.8	Management Summary	81
5	Härtung.....	83
5.1	Dedizierte KI-Hosts und Linux-Baselines	84
5.2	MAC und Laufzeitschutz.....	87
5.3	Container- und Kubernetes-Härtung.....	89

5.4	GPU- und Ressourcen-Zugriffe	91
5.5	GPU-Isolation und MIG-Sicherheit	93
5.6	Service-Accounts und Rechte-Minimierung	95
5.7	Admission- und Policy-Mechanismen	97
5.8	Hardware Root of Trust und Confidential Computing	99
5.9	Model Supply Chain Security	103
5.10	Management Summary	104
6	Segmentierung und Service-Trust	105
6.1	Netzwerkzonen für KI-Plattformen	106
6.2	Default-Deny und Mikrosegmentierung	109
6.3	Kubernetes Network Policies in der Praxis	111
6.4	mTLS und Service-Mesh-Ansätze	113
6.5	Egress-Kontrolle und Exfiltrationsschutz	117
6.6	Häufige Netzwerkfehlkonfigurationen	119
6.7	Management Summary	121
7	Absicherung der RAG- und Datenpipeline	123
7.1	Risiken in der Dokument-Ingestion	124
7.2	Dateiprüfung, Malware-Scan und CDR	127
7.3	CDR-Werkzeuge	130
7.4	Dokumentversionierung und Nachvollziehbarkeit	131
7.5	Zugriffskontrolle im Retrieval	133

7.6	Mandantentrennung im Vektorstore.....	137
7.7	Angriffsszenarien und Gegenmaßnahmen	141
7.8	PII-Erkennung und Datenschutzsäuberung	143
7.9	Management Summary	147
8	Prompt-Security und Guardrails.....	149
8.1	Anatomie von Prompt-Injection-Angriffen	150
8.2	Trennung von System-, Nutzer- und Kontextdaten.....	154
8.3	Heuristische und modellbasierte Filter	156
8.4	Guardrails-Frameworks im Einsatz	158
8.5	Output-Validierung und Policy Enforcement	161
8.6	Defense-in-Depth für LLM-Anwendungen.....	164
8.7	Agentic AI und Tool-Calling-Sicherheit	166
8.8	Sicherheitsanforderungen für Agentic-AI-Systeme	169
8.9	Semantic Web Application Firewall (sWAF).....	172
8.10	Management Summary	176
9	Monitoring, Detection und Betrieb	177
9.1	Observability für GenAI-Plattformen	178
9.2	Wichtige Betriebs- und Sicherheitsmetriken.....	183
9.3	Prompt- und Nutzungsanomalien erkennen	185
9.4	Drift-, Abuse- und Token-Flood-Detection.....	187
9.5	Runbooks für Updates und IR.....	191

9.6	Model Lifecycle Management	197
9.7	Kontinuierliche Sicherheitsüberprüfung	200
9.8	Management Summary	203
10	Governance, Compliance und Ausblick	205
10.1	Regulatorische Anforderungen	206
10.2	Datenklassifizierung und Audit-Trails.....	211
10.3	Organisatorische Betriebsmodelle.....	216
10.4	Typische Fehlkonfigurationen in Audits.....	219
10.5	Technologische Trends bei lokaler KI.....	222
10.6	Handlungsempfehlungen	224
10.7	Reifegradmodell für GenAI-Plattformen	229
10.8	Management Summary	230
	Zum Schluss	231
	Anhang.....	VII
	Checkliste: Secure Private AI Readiness	VII
	Checkliste: GenAI-Architekturentscheidungen	XII
	Referenz-Threat-Model für RAG-Systeme	XVI
	Abkürzungverzeichnis	XXV
	Literatur- und Quellenverzeichnis	XXIX
	Stichwortverzeichnis	XXXIII

Mehr von Brain-Media.de	XLI
IT-Texter.one	XLIV

Vorwort

Generative KI hat sich in erstaunlich kurzer Zeit von einem experimentellen Forschungsthema zu einer produktiven Infrastrukturkomponente entwickelt. Große Sprachmodelle beantworten Supportanfragen, analysieren interne Dokumente, unterstützen Entwickler und automatisieren Wissensarbeit. Parallel dazu wächst in vielen Organisationen der Wunsch, diese Fähigkeiten nicht ausschließlich über externe Cloud-Dienste zu nutzen, sondern unter eigener Kontrolle zu betreiben. Lokale oder privat betriebene GenAI-Plattformen versprechen Datenhoheit, geringere Latenzen und bessere Integrationsmöglichkeiten in bestehende IT-Landschaften. Doch mit der technischen Kontrolle wandert auch die volle Verantwortung für Sicherheit, Stabilität und Governance ins eigene Haus.

Viele erste Implementierungen folgen noch einem prototypischen Muster: Modell laden, API starten, Chat-Frontend anbinden – und produktiv gehen. Was in der Evaluierungsphase funktioniert, offenbart im Dauerbetrieb jedoch schnell strukturelle Schwächen. Große Sprachmodelle bringen eine eigene Angriffsdynamik mit. Prompts können manipuliert werden, Retrieval-Bestände lassen sich vergiften, und ungeschützte Inferenzendpunkte eröffnen neue Missbrauchsszenarien. Gleichzeitig verbinden moderne GenAI-Stacks mehrere bislang getrennte Disziplinen: GPU-beschleunigte Inferenz, datengetriebene Retrieval-Pipelines, API-Gateways, Container-Orchestrierung

und oft noch unreife Entwickler-Workflows. Risiken entstehen dabei selten durch einen einzelnen gravierenden Fehler, sondern durch unsaubere Übergänge zwischen diesen Komponenten.

Hinzu kommt eine begriffliche Unschärfe. In der Forschung bezeichnet „Private AI“ – ich spreche in diesem Buch von lokaler KI – häufig Verfahren wie homomorphe Verschlüsselung, Differential Privacy oder föderiertes Lernen. In der betrieblichen Praxis hingegen geht es meist um lokal oder in kontrollierten Umgebungen ausgeführte GenAI-Workloads. Dieses Buch folgt bewusst der zweiten Perspektive: Im Mittelpunkt steht der sichere Entwurf, die Härtung und der stabile Betrieb lokal kontrollierter GenAI- und RAG-Systeme. Kryptographische Spezialverfahren werden dort betrachtet, wo sie architektonisch relevant sind, nicht als alleinige Lösung für KI-Sicherheit.

Der Schwerpunkt liegt klar auf umsetzbaren Maßnahmen. Statt abstrakter Grundsatzdiskussionen behandelt das Buch konkrete Architekturentscheidungen, typische Fehlkonfigurationen und praxiserprobte Schutzmechanismen. Anhand einer durchgängigen Referenzarchitektur wird gezeigt, wie sich eine lokale KI-Plattform strukturiert aufbauen und absichern lässt – von der Netzwerksegmentierung über die Absicherung der RAG-Pipeline bis hin zu Guardrails, Monitoring und Incident-Runbooks. Ziel ist eine Umgebung, die nicht nur funktional überzeugt, sondern auch unter Sicherheits-, Compliance- und Betriebsaspekten beherrschbar bleibt.

Das Buch richtet sich an Architekten, Plattform- und DevOps-Teams, Security Engineers sowie technisch orientierte Entscheider, die GenAI-Workloads im eigenen Verantwortungsbereich betreiben oder dies planen. Vorausgesetzt werden Grundkenntnisse in Container- und Cloud-nativen Umgebungen; tiefes Vorwissen zu großen Sprachmodellen ist hilfreich, aber nicht zwingend erforderlich.

Lokale KI wird in den kommenden Jahren fester Bestandteil moderner Unternehmensinfrastrukturen werden. Organisationen, die früh eine saubere Architekturdiziplin, klare Trust Boundaries und belastbare Betriebsroutinen etablieren, verschaffen sich einen nachhaltigen Vorteil.

Dieses Buch soll dabei als technischer Leitfaden dienen — mit dem Anspruch, aus schnellen Proof-of-Concepts dauerhaft kontrollierbare, sichere KI-Plattformen zu machen.

Ich wünsche Ihnen viel Erfolg beim Betrieb und der Absicherung lokaler KI-Systeme.

Herzlichst

Holger Reibold

1 Private KI im Infrastrukturkontext

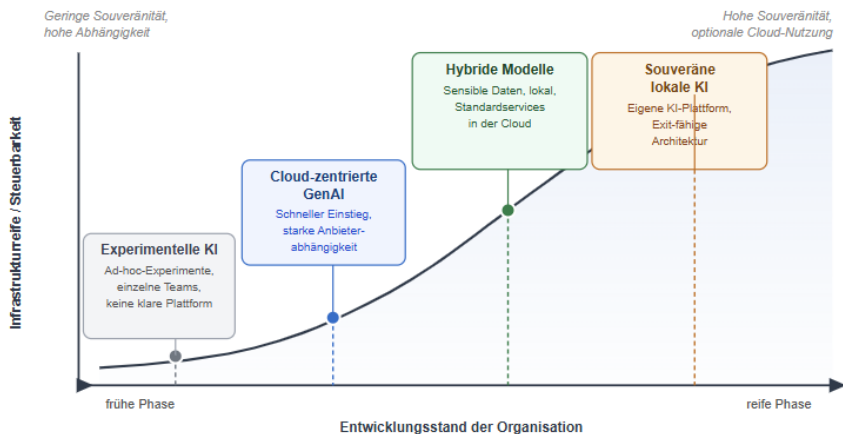
Lokale GenAI ist dort alternativlos, wo Datenabfluss regulatorisch oder organisatorisch nicht tolerierbar ist.

Lokale GenAI-Systeme entwickeln sich zunehmend von isolierten Pilotprojekten zu festen Bestandteilen der Unternehmensinfrastruktur. Während frühe Experimente meist mit minimalem Betriebs- und Sicherheitsfokus durchgeführt wurden, erfordern produktive Deployments eine deutlich höhere architektonische Disziplin. Modelle, Retrieval-Pipelines und Orchestrierungsschichten müssen heute wie kritische Plattformkomponenten behandelt werden – mit klaren Trust Boundaries, reproduzierbaren Deployments und überprüfbaren Datenflüssen.

Dieses Kapitel ordnet lokale beziehungsweise privat betriebene KI in den Infrastrukturkontext moderner IT-Landschaften ein. Im Mittelpunkt stehen die Treiber für On-Prem-GenAI, typische Einsatzmuster sowie die zentralen Architekturentscheidungen zwischen Cloud- und lokalem Betrieb. Ziel ist es, eine belastbare Entscheidungsgrundlage zu schaffen, bevor in den folgenden Kapiteln die konkrete Referenzarchitektur, Härtingsmaßnahmen und Betriebsprozesse vertieft werden. Wer die infrastrukturellen Rahmenbedingungen richtig

bewertet, vermeidet viele der späteren Sicherheits- und Skalierungsprobleme bereits in der Entwurfsphase.

Die in nachstehender Abbildung skizzierte Entwicklung zeigt den typischen Reifungspfad von KI-Infrastrukturen in Organisationen. Viele Initiativen beginnen in einer experimentellen Phase, die durch schnelle Cloud-Nutzung, geringe Governance und hohe Flexibilität geprägt ist. In diesem Stadium steht die funktionale Machbarkeit im Vordergrund, während Aspekte wie Kostenstabilität, Datenhoheit oder Betriebsreife noch eine untergeordnete Rolle spielen.



Typischer Entwicklungspfad von KI-Infrastrukturen von experimentellen Cloud-Nutzungen hin zu souveränen, produktionsreifen Plattformen mit zunehmender Governance-, Kosten- und Kontrolltiefe.

Mit wachsender Nutzung und zunehmender geschäftlicher Relevanz verschiebt sich der Fokus jedoch deutlich. Organisationen professionalisieren ihre Plattformen, standardisieren Betriebsmodelle und beginnen, regulatorische sowie ökonomische Effekte systematisch zu bewerten. In dieser Phase entstehen häufig hybride Architekturen, die Cloud-Agilität mit ersten lokalen Komponenten kombinieren.

Die rechte Seite des Entwicklungspfads markiert schließlich die Phase souveräner KI-Plattformen. Hier stehen kontrollierte Datenflüsse, reproduzierbarer Betrieb, Kosten-transparenz und technische Durchsetzbarkeit von Sicherheitsanforderungen im Mittelpunkt. Lokale oder dediziert betriebene GenAI-Workloads sind in diesem Stadium keine Gegenbewegung zur Cloud, sondern Ausdruck architektonischer Reife und bewusster Workload-Platzierung.

Hier zeigt sich: Der Übergang zu lokaler oder hybrider GenAI ist in vielen Organisationen kein disruptiver Strategiewechsel, sondern die konsequente Weiterentwicklung wachsender KI-Produktionsumgebungen.

1.1 Vom Experiment zur produktiven Plattform

Der Einsatz generativer KI beginnt in den meisten Organisationen experimentell. Einzelne Teams evaluieren Modelle, testen Chat-Oberflächen oder integrieren erste Retrieval-Funktionen in bestehende Anwendungen. Diese Phase ist bewusst explorativ: Geschwindigkeit zählt mehr als Robustheit, Sicherheitskontrollen sind minimal, und

die Infrastruktur ist häufig ad hoc aufgebaut. Typisch sind lokal gestartete Modellserver, einfache API-Gateways und unstrukturierte Dokumentenbestände für RAG-Tests.

Mit wachsender Nutzung verschieben sich jedoch die Anforderungen grundlegend. Sobald GenAI-Funktionen in produktive Workflows eingebunden werden – etwa im Support, in der Softwareentwicklung oder in internen Wissenssystemen – entstehen klassische Plattformanforderungen: Verfügbarkeit, Skalierbarkeit, Mandantentrennung, Nachvollziehbarkeit und Betriebssicherheit. Ein Sprachmodell ist dann nicht mehr nur ein Werkzeug, sondern Bestandteil einer geschäftskritischen Verarbeitungskette.

Der zentrale Fehler vieler Organisationen besteht darin, Proof-of-Concept-Architekturen nahezu unverändert in den Produktivbetrieb zu überführen. Dabei fehlen meist mehrere Schichten, die für einen stabilen Plattformbetrieb erforderlich sind: reproduzierbare Build- und Deployment-Pipelines, versionierte Modellartefakte, klar definierte Service-Identitäten, kontrollierte Netzwerkzonen und belastbares Monitoring. Besonders kritisch ist die Datenebene. In frühen Tests werden Dokumente häufig ungeprüft ingestiert und global durchsuchbar gemacht. Unter Produktionsbedingungen führt dies schnell zu Compliance- und Sicherheitsproblemen.

Ein weiterer Unterschied zwischen Experiment und Plattformbetrieb liegt in der Deterministik. Klassische Microservices liefern bei gleichen Eingaben reproduzierbare Ergebnisse. Große Sprachmodelle hingegen erzeugen probabilistische Antworten, deren Qualität stark

vom Prompt, vom Kontext und von der Modellversion abhängt. Daraus folgt, dass Konfigurationsänderungen – etwa neue Embeddings, geänderte Chunk-Größen oder aktualisierte Modellgewichte – unmittelbare Auswirkungen auf das Laufzeitverhalten haben können. Produktive Plattformen benötigen deshalb ein sauberes Versionierungs- und Rollback-Modell nicht nur für Code, sondern auch für Modelle und RAG-Datenbestände.

Auch die Bedrohungslage verändert sich mit zunehmender Exposition. Während interne Tests meist von vertrauenswürdigen Nutzern durchgeführt werden, sind produktive Systeme systematisch automatisierten Abfragen, Prompt-Injection-Versuchen und möglichen Datenabfluss-Szenarien ausgesetzt. Ohne Rate Limits, Eingabevalidierung, kontextbewusste Prompt-Trennung und mehrstufige Guardrails steigt das Risiko exponentiell mit der Nutzung.

Der Übergang zur produktiven GenAI-Plattform ist daher primär eine Frage der Architekturdisziplin. Erfolgreiche Umgebungen zeichnen sich durch klar getrennte Sicherheitszonen, identitätsbasierte Dienstkommunikation, streng kontrollierte Ingestion-Pipelines und durchgängige Observability aus. Erst wenn diese Grundlagen etabliert sind, lässt sich GenAI unter realer Last zuverlässig betreiben. Organisationen, die diesen Reifegrad systematisch erreichen, verwandeln experimentelle KI-Funktionen in eine dauerhaft beherrschbare Infrastrukturkomponente.

1.2 Treiber für On-Prem-GenAI

Der Trend zu lokal betriebenen GenAI-Systemen ist kein Selbstzweck, sondern die Folge konkreter technischer und regulatorischer Anforderungen. Drei Faktoren dominieren die Architekturentscheidungen in der Praxis: Datenkontrolle, Latenzverhalten und Compliance-Fähigkeit. Ihre Gewichtung variiert je nach Branche, bestimmt jedoch maßgeblich, ob Cloud-basierte LLM-Dienste ausreichen oder eine On-Prem-Plattform notwendig wird.

Der wichtigste Treiber ist die Datenhoheit. Viele produktive GenAI-Anwendungen greifen auf interne Wissensbestände zu – Quellcode, Verträge, Forschungsdaten, Supporttickets oder personenbezogene Informationen. In solchen Szenarien reicht eine vertragliche Zusicherung externer Anbieter oft nicht aus. Organisationen benötigen technisch überprüfbare Garantien darüber, wo Daten verarbeitet werden, welche Telemetrie abfließt und ob Inhalte in Trainingspipelines einfließen können. Lokale Deployments ermöglichen eine klare Kontrolle der Datenpfade, insbesondere wenn Modellendpunkte in isolierten Netzwerkzonen betrieben und ausgehende Verbindungen restriktiv gehandhabt werden. Für sicherheitskritische Umgebungen ist diese physische und logische Kontrollierbarkeit ein entscheidender Vorteil.

Die in nachstehender Abbildung dargestellten Einflussfaktoren wirken in der Praxis selten isoliert, sondern verstärken sich häufig gegenseitig. In vielen Organisationen ist beispielsweise nicht allein der

Datenschutz ausschlaggebend, sondern die Kombination aus regulatorischem Druck, steigenden Inferenzkosten und wachsenden Anforderungen an Betriebskontrolle. Das Kräftebild verdeutlicht damit, dass die Entscheidung für On-Prem-GenAI typischerweise aus einem mehrdimensionalen Bewertungsprozess hervorgeht.



Die zentralen Treiber für den Einsatz von On-Prem-GenAI. Die relative Gewichtung von Datenschutz, Latenz, Kosten, Kontrolle und Compliance bestimmt die optimale Platzierung von GenAI-Workloads.

Auffällig ist zudem, dass die Gewichtung der einzelnen Treiber stark vom Anwendungskontext abhängt. Während in hochregulierten Branchen Datenschutz und Compliance dominieren, stehen in interaktiven oder produktionsnahen Szenarien häufig Latenz und deterministische Performance im Vordergrund. Kosten entfalten ihre Wirkung oft zeitverzögert: Mit wachsendem Inferenzvolumen verschiebt sich die Wirtschaftlichkeitsbetrachtung nicht selten zugunsten lokaler Plattformen.

Der zweite wesentliche Faktor ist die Latenz. GenAI-Workloads, insbesondere RAG-Systeme mit großen Kontextfenstern, generieren hohe Datenvolumina pro Anfrage. Jede externe Round-Trip-Kommunikation erhöht die Antwortzeit und verschlechtert die Interaktivität. In Entwicklungsumgebungen, internen Assistenten oder automatisierten Workflows kann bereits eine zusätzliche Sekunde spürbare Produktivitätseinbußen verursachen. Lokale Inferenz reduziert Netzwerkabhängigkeiten und stabilisiert Antwortzeiten auch bei hoher Parallelität. Dieser Effekt wird mit zunehmender Tokenzahl, komplexeren Retrieval-Schritten und multimodalen Modellen weiter verstärkt.

Der dritte Treiber ist regulatorischer Druck. Datenschutzgesetze, branchenspezifische Vorgaben und interne Governance-Richtlinien verlangen nachvollziehbare Datenverarbeitung, feingranulare Zugriffskontrollen und auditierbare Systemzustände. Lokale GenAI-Plattformen erleichtern die Umsetzung dieser Anforderungen erheblich. Dokumentbasierte Zugriffskontrollen im RAG-Store,

Literatur- und Quellenverzeichnis

- Abadi, M., et al. (2016). Deep Learning with Differential Privacy. Advances in Neural Information Processing Systems, 35, 30318–30332. <https://arxiv.org/abs/2208.07339>
- AMD (2025). The Rise of AI PCs in the Commercial Market: Planning for the Future. <https://www.amd.com/en/blogs/2025/the-rise-of-ai-pcs-in-the-commercial-market-plann.html>
- Anthropic. (2022). Constitutional AI: Harmlessness from AI feedback. Online: <https://www.anthropic.com>.
- Apple Inc. (2023). MLX: Efficient and flexible machine learning on Apple.
- Arize AI. (2023). Phoenix: Open-source AI observability for LLMs and RAG.
- Barnett, S., et al. (2024). Seven Failure Points When Engineering a Retrieval Augmented Generation System.
- Cloud Native Computing Foundation. (2023). Kubernetes security best practices. <https://kubernetes.io>
- Dettmers, T., et al. (2023). QLoRA: Efficient Finetuning of Quantized LLMs.
- Dettmers, T., Lewis, M., Belkada, Y., & Zettlemoyer, L. (2022).
- Dubey, A., et al. (2024). The Llama 3 herd of models. arXiv. <https://arxiv.org/abs/2407.21783>
- European Parliament and Council of the European Union. (2016). Regulation (EU) 2016/679 (General Data Protection Regulation). Online. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32016R0679>
- European Parliament and Council of the European Union. (2024). Regulation (EU) 2024/1689 (Artificial Intelligence Act). Online:

<https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689>.

- Frantar, E., Ashkboos, S., Hoefler, T., & Alistarh, D. (2023). GPTQ: Accurate post-training quantization for generative pre-trained transformers. International Conference on Learning Representations.
- Gerganov, G. (2023). llama.cpp [Computer software]. GitHub. <https://github.com/ggerganov/llama.cpp>
- Google. (2023). Secure AI Framework (SAIF). Google Cloud. <https://cloud.google.com/security/saif>
- Han, S., Mao, H., & Dally, W. J. (2015). Deep compression: Compressing deep neural networks with pruning, trained quantization and Huffman coding.
- Hu, E. J., et al. (2021). LoRA: Low-Rank Adaptation of Large Language Models.
- Inan, H., et al. (2023). Llama Guard: LLM-based Input-Output Safeguard for Human-AI Conversations.
- Jiang, A. Q., et al. (2023). Mistral 7B. arXiv:2310.06825.
- Johnson, J., et al. (2019). Billion-scale similarity search with GPUs.
- Khowaja, S. A., Dev, K., Qureshi, N. M. F., Khuwaja, P., & Foschini, L. (2022). Toward industrial private AI: A two-tier framework for data and model security. IEEE Wireless Communications, 29(2), 76–83.
- Kwon, W., et al. (2023). Efficient memory management for large language model serving with PagedAttention. Proceedings of the ACM Symposium on Operating Systems Principles, 611–626.
- Lauter, K. (2022). Private AI: machine learning on encrypted data. In Recent advances in industrial and applied mathematics (pp. 97–113). Cham: Springer International Publishing.
- Lewis, P., Perez, E., Piktus, A., et al. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. Advances in Neural

- Information Processing Systems, 33, 9459–9474.
<https://arxiv.org/abs/2005.11401>
- Lin, J., Tang, J., Tang, H., Yang, S., Dang, X., & Han, S. (2023). AWO: Activation-aware weight quantization for LLM compression and acceleration. arXiv. <https://arxiv.org/abs/2306.00978>
- LocalAI Project. (2023). LocalAI: The free, Open Source OpenAI alternative [Computer software]. <https://localai.io>
- McMahan, B., et al. (2017). Communication-Efficient Learning of Deep Networks from Decentralized Data.
- Microsoft Corporation. (2021). ONNX Runtime. Online: <https://onnxruntime.ai>.
- MITRE. (2023). Adversarial Threat Landscape for Artificial-Intelligence Systems (ATLAS). <https://atlas.mitre.org>
- NIST. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). <https://www.nist.gov>
- NVIDIA Corporation. (2023). TensorRT-LLM. Online: <https://github.com/NVIDIA/TensorRT-LLM>
- Ollama Team. (2023). Ollama [Computer software]. <https://ollama.com>
- Open Source Initiative (OSI). (2024). The Open Source AI Definition.
- OWASP Foundation. (2023). OWASP Top 10 for large language model applications. <https://owasp.org>.
- Rebedea, T., et al. (2023). NeMo Guardrails: A Toolkit for Steerable and Safe LLM Applications. (NVIDIA).
- Vaswani, A., et al. (2017). Attention is all you need. Advances in Neural Information Processing Systems, 30.
- von Platen, P., et al. (2022). Diffusers: State-of-the-art diffusion models. GitHub.
- Weidinger, L., et al. (2021). Ethical and social risks of harm from language models.

- Wolf, T., et al. (2020). Transformers: State-of-the-art natural language processing. *EMNLP System Demonstrations*, 38–45.
- Xiao, G., Lin, J., Seznec, M., Wu, H., Demouth, J., & Han, S. (2023). SmoothQuant: Accurate and efficient post-training quantization for large language models. *International Conference on Machine Learning*.
- Zhou, Z., Chen, X., Li, E., Zeng, L., Luo, K., & Zhang, J. (2019). Edge intelligence: Paving the last mile of artificial intelligence with edge computing. *Proceedings of the IEEE*, 107(8), 1738–1762.

Stichwortverzeichnis

A

Abuse-Muster.....	37, 68
Admission Controller	97
Agentic AI.....	166
AI RMF	40
Air-Gapped-Deployment	63, 79
Allowlist.....	56
AMD SEV	76
Analyse.....	14
Anbietern.....	20
Angriffsfläche.....	25
Angriffsmuster	29, 141
Anomalieerkennung	28
Anwendungssicherheit	29
Apache Tika.....	130
API.....	1
Application Zone.....	46
Applikationszone	106
Architekturebene	43
Architekturentscheidung	15
Artefaktimport.....	79
Assistenzsystem	14
ATLAS.....	49
Auditierbarkeit	56

Audit-Logging	85
Authentisierung.....	85
Automatisierung.....	14
Automounting.....	96
AWQ	222

B

Batch Inference	21
Bedrohungslage	184
Bedrohungsmodell.....	23
Beobachtbarkeit.....	47
Bootpfad.....	99

C

CDR.....	34
CDR-Werkzeuge	130
CIDR.....	120
Cloud	15, 18
ClusterRoles	90
Codeanalyse.....	14
Compliance-by-Design	208
Compliance-Matrix	208
Compute-Pool	84
Confidential Computing	101

Container-Image	90
Containerisierung	89
Container-Orchestrierung	1
Content Disarm and Reconstruction	130

D

Data Poisoning	33
Dateiprüfung	127
Datenebene	66
Datenexfiltration	14, 151
Datenfluss.....	77
Datenhoheit.....	10
Datenhygiene	201
Datenklassifizierung	211
Datenpfad.....	123
Datenpipeline.....	123
Datenschutzrecht	206
Datenvergiftung	33
Datenzone.....	107
Default-Deny	109
Default-Deny-Prinzip.....	43
Defense-in-Depth	164
Denial-of-Service	92
Deployment.....	5
Designprinzip	44
Dienstkommunikation	55
Dokument-Ingestion	72

Dokumentversion	132
DSGVO	123, 143, 206

E

Egress.....	39, 105
Egress-Kontrolle	117
Einsatzszenario	13
Embedding.....	72
Embedding-Dienst.....	75
Enforce	98
EU AI Act	207
Exfiltrationsschutz	117
ExifTool	130
Extraktion	129

F

False Negatives.....	145
False Positives	145
Fehlverhalten	30
Filter.....	157
Fine-Tuning.....	35
Fine-Tuning-Modell	35

G

Gateway	26, 68
GenAI.....	1
GenAI-Workload	7

Generative KI.....	22
Governance-Strategie	209
GPTQ	222
GPU.....	15
GPU-Cluster.....	19
GPU-Isolation.....	93
GPU-Server	20
Guardrails.....	2
Guardrails-Framework	158

H

Hardening.....	51
Hardware Root of Trust.....	99
Härtung.....	83
Heuristik	156

I

Incident-Runbook.....	2
Inference-Zone	63
Inferenz.....	12
Inferenzzone	107
Ingestion-Worker	50
Input Validation.....	29
Intel TDX.....	76
IoT	15
IP-Allowlisting	55
Isolation	76
IT 1	

IT-Security.....	19
------------------	----

J

Jailbreak.....	32
JSON	161

K

KI-Host.....	84
KI-Plattform	2
Klassifikationsscore.....	157
Klassifikator.....	69
Kopplungstiefe.....	16
KRITIS	62
Kubernetes.....	50
Kubernetes Network Policies..	111
Kubernetes-Cluster.....	105

L

Latenz	12
Laufzeitschutz	88
Linux-Baseline	84
LLM.....	15
LLM-Orchestrator	53
Logging	39
Lokale KI.....	2
LoRA	35
Low-Rank Adaptation.....	35

M

MAC	87
Malware-Scan	74, 128
Mandantentrennung.....	47, 137
Mandatory Access Control.....	87
Masking	144
Massennutzung.....	37
MCP	170
Mikrosegmentierung	108
MITRE-ATLAS	39
ML-Engineering	19
Model Context Protocol.....	170
Model Extraction.....	28, 37, 39
Model Lifecycle Management	197
Model Serving.....	26, 56, 75
Model Supply Chain Security	103
Modellabfrage.....	23
Modellartefakt.....	66
Modellhosting	15
Modellinteraktion.....	46
Modelllebenszyklus	20
Modellrekonstruktion	28
Monitoring.....	2, 178
Monorepo.....	14
mTLS	49, 113
Multi-Instance GPU.....	92
Mutual Transport Layer Security	114

N

Network Policies	90
Netzwerkfehlkonfiguration.....	119
Netzwerksegmentierung.....	105
Netzwerkzone.....	106
NIST	40
Nutzungsanomalie	185
Nutzungsmetrik	183

O

Observability.....	9
OIDC/OAuth-Flow	68
On-Prem	18
On-Prem-GenAI	5, 10
OPSWAT	130
Orchestrierung.....	26, 51
Orchestrierungsschicht .	43, 63, 70
Output-Validierung.....	74, 161
OWASP Top 10	40

P

Patchmanagement	15
Performance.....	85
Perimeter-Schutzmodell	23
PII	126
PII-Erkennung.....	143
Pipeline.....	194

Plattformbetrieb.....	19
Pod.....	91
Policy Enforcement.....	74, 161
Policy Engines.....	91
Policy-as-Code	113
Private AI	2
Prompt Injection	23, 31
Prompt-Anomalie	185
Prompt-Architektur	32
Promptgröße.....	178
Prompt-Injection-Angriff	150
Prompt-Konkatenation	154
Prompt-Logik.....	51
Prompt-Manipulation	33
Prompt-Security.....	149
Prompt-Verarbeitung.....	149
Proof-of-Concept.....	8
Proxy.....	43
Pseudo-Default-Deny.....	110

Q

QLoRA	35, 222
Quellcode-Repository	35
Query	37
Quotensteuerung.....	68

R

RAG.....	8
----------	---

RAG-Datenpflege.....	20
RAG-Pipeline.....	72
RAG-Store	12
Rate Limit	37
Rate Limiting.....	68
Rebuild-Zyklus	85
Red Teaming.....	201
Redaction	144
Refactoring.....	14
Referenzarchitektur	61
Referenzmuster.....	57
Regex.....	29
Reifegradmodell	229
Reifungspfad.....	6
Reindexierung.....	146
Repository-Navigation	14
Ressourcensteuerung	92
Retrieval-Angriff	34
Retrieval-Augmented Generation	13
Retrieval-Manipulation.....	23
Retrieval-Steuerung	51
Review-Zyklus.....	202
Risiko-Heatmap.....	25
Role-Based Access Control	90
Rollback.....	9
Runbook	191

S

SaaS.....	80
Sasa Software	130
Schema-Validierung	29
Schichtenarchitektur.....	87
Security Context	90
Segmentierung	45
Semantic Web Application	
Firewall	149, 172
Service-Account.....	95
Sicherheitsmetrik.....	183
Sicherheitsüberprüfung	200
Sicherheitsziel	44
Sicherheitsziele	84
Sidecar-Pattern	52
SIEM	38
Signaturprüfung	36
Skalierung	15
Soft Delete.....	146
Software Engineering.....	14
Sprachmodell.....	14
STRIDE.....	39
sWAF	32, 58

T

Telemetrie	64
Telemetriemodell.....	179

Threat Model.....	40
Token Flooding	37
Token Harvesting.....	37
Token-Flooding.....	188
Tokenverbrauch.....	178
Tool-Call Escalation.....	168
Toolchain	222
TPM	99
Traceability	132
Trainingsumgebungen.....	36
Treiber	10
Trend	222
Trust Boundaries	30
Trust Boundary	47
Trusted Platform Module.....	99

U

Überprovisionierung	77
---------------------------	----

V

Vektorstore.....	28, 65
Vertraulichkeit.....	44
Votiro	130
VRAM	183

W

Wirtschaftlichkeit.....	19
-------------------------	----

Wissensspeicher 66

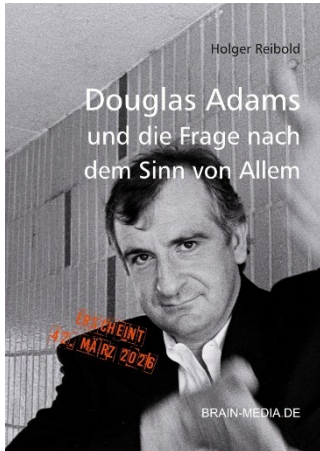
X

X.509-Zertifikat 55

Z

Zonenmodell 106

Zugriffskontrolle 133



42 – Douglas Adams und die Frage nach dem Sinn von Allem

Am 11. Mai 2026 ist Douglas Adams 25 Jahre tot. Der Kultautor hat der Welt wunderbar, skurrile Werke geschenkt. Jetzt ist es an der Zeit, den Autor kennenzulernen.

Umfang: 140 Seiten

Preis: 14,99 EUR

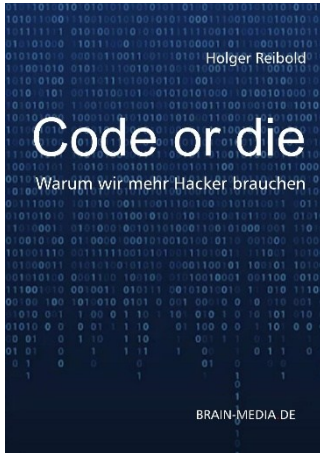
Erscheint: 42. März 2026



Towelday, das ultimative Handtuch für alle Fans

An seinem Todestag, dem Towelday, erinnern sich Fans an Douglas Adams und huldigen dem Kultautor.

100 % intergalaktisch geprüfte Baumwolle, nachhaltig Produktion zum Preis von 42 EUR.



Code or die – Warum wir mehr Hacker brauchen

Ein Manifest für mehr digitale Selbstbestimmung, Neugierde und Eigenverantwortung. Medienkompetenzen alleine genügen nicht; die Gesellschaft von morgen braucht Digitalkompetenzen.

Umfang: 120 Seiten

Preis: 14,99 EUR

Erscheint Frühjahr 2026



KI Incident Response – Wie man Sicherheitsvorfälle in KI-Systemen erkennt, eindämmt und verantwortet

Ziel- und punktgenaue Reaktionen für kritischen KI-Vorfälle.

Umfang: 140 Seiten

Preis: 16,99 EUR

Erscheint: Frühjahr 2026



Brain-Media KaaS

Individual

Business

Enterprise

– die E-Book-Flatrate –

Brain-Media.de ist ein **Pionier** in Sachen E-Book-Flatrates. Bereits 2014 wurden mit Plus+ E-Books als attraktives und preisgünstiges Abo angeboten. Im April 2026 erlebt Plus+ eine Renaissance.

Abonnenten erhalten **1 Jahr Zugriff auf alle E-Books**, die bereits erschienen sind, und auf alle Neuerscheinungen des Abozeitraums. Sie können über 30 verfügbare E-Books herunterladen – dazu editierbare Checklisten, Vorlagen, Beispiele etc. Sie können die E-Books ausdrucken und auf mehreren Lesegeräten verwenden. Die E-Books verwenden kein DRM!

Und das alles zum **unschlagbaren Sonderpreis!**

IT-Texter.one

100+ IT-Fachbücher

1500+ Fachartikel

30+ Erfahrung

KOMPLEXE INHALTE PUNKTGENAU AUFZUBEREITEN, IST EINE KUNST. ICH BEHERRSCHE SIE. BEI MIR ERHALTEN SIE FACH-
TEXTE, DIE KOMPLEXES VERSTÄNDLICH MACHEN.

Seit über 30 Jahren unterstütze ich Unternehmen aus der IT-, Software- und Digitalbranche dabei, ihre technischen Inhalte klar, präzise und zielgruppenorientiert zu kommunizieren. Als promovierter Informatiker und erfahrener IT-Journalist verbinde ich fundiertes Fachwissen mit journalistischem Storytelling. Als Key Account Manager eines IT-Dienstleisters verfüge ich obendrein über konkrete Erfahrungen mit allen gängigen Technologien.

WARUM SIE MIT MIR ARBEITEN SOLLTEN

35 Jahre Erfahrung mit Internet-,
Netzwerk- und Webtechnologien

Kooperation mit führenden Akteuren
der IT- und Medienbranche

Strategisches Denken: Texte, die nicht nur informieren,
sondern auch verkaufen

THEMENSCHWERPUNKTE	WIE KANN ICH SIE UNTERSTÜTZEN?
Open-Source	Content Creation
Enterprise IT	Dokumentationen
IT-Consulting	Case Studies
SaaS	Suchmaschinenoptimierung
Künstliche Intelligenz	Tech-Marketing

MEIN VERSPRECHEN

Ich übernehme die inhaltliche und sprachliche Brücke zwischen Technologie und Anwendung. Selbst komplexe Sachverhalte kommen beim Publikum an – fachlich korrekt, prägnant und SEO-wirksam.

PREISMODELLE

Professionelle Leistungen, die ihresgleichen suchen, gibt es nicht umsonst. Sprechen Sie mit an. Gerne vereinbaren wir einen Fixpreis; das vereinfacht Ihre Kalkulation.

KONTAKT AUFNEHMEN

Sprechen wir über Ihr Projekt. Schreiben Sie mir eine Mail (info@it-texter.one). Oder besser noch: Rufen Sie mich an (+49 681 91005698).