

AI-Operable IT-Services

Strategisches Whitepaper zur Gestaltung autonom
betreibbarer IT-Anwendungen und IT-Services

*Wie das Zusammenspiel aus AI-operablem Service, Governance,
Agentensystem und AI Service Brain autonome IT-Services ermöglicht*

Autor: Markus Berg

Organisation: Hessischer Rundfunk (hr), Mitglied der ARD

Version: 3.7

Stand: 16. Juni 2026

Status: Konsolidierte Machbarkeitsstudie / Forschungsprogramm

Diese Studie entsteht im Hessischen Rundfunk (hr) als Mitglied der ARD. Sie untersucht, wie KI-Agentensysteme den Betrieb und die Weiterentwicklung von IT-Services des öffentlich-rechtlichen Rundfunks künftig unterstützen können.

Inhaltsverzeichnis

Inhaltsverzeichnis	2
Abbildungsverzeichnis	5
1. Executive Summary	6
2. Ausgangssituation und Motivation	7
2.1 Der aktuelle Stand: KI als Werkzeug	7
2.2 Die strukturelle Grenze	7
2.3 Der Perspektivwechsel	7
2.4 Warum jetzt	7
3. Vision und Zielbild	9
3.1 Vision	9
3.2 Verantwortungsteilung im Zielbild	9
3.3 Die drei parallel entstehenden Artefakte	10
4. Zentrale Forschungsfragen und Hypothesen	11
4.1 Leitende Forschungsfrage	11
4.2 Arbeitshypothesen	11
4.3 Übergeordnete Forschungsfrage zur Selbst-Evolution	11
5. Die sechs Perspektiven der Studie	12
5.1 Betriebsautomatisierung	12
5.2 Autonomer Servicebetrieb	12
5.3 Autonome Service-Evolution	12
5.4 AI-Operability by Design	13
5.5 Agent Engineering	13
5.6 Governance & Trust	14
6. Referenzmodell für AI-Operable IT-Services	16
6.1 Grundannahme	16
6.2 Die fünf Ebenen	16
6.3 Der Agent Layer im Detail	18
6.4 Referenzarchitektur des Demonstrators	18
6.5 AI-Operability by Design — die sechs Prinzipien	19
7. AI-Readiness-Modell	21
7.1 Definition	21
7.2 Die acht Dimensionen	21
7.3 Bewertung und Reifegrad	22

8. Governance- und Trust-Modell	24
8.1 Die Leitfrage der Governance.....	24
8.2 Grundprinzipien.....	24
8.3 Die fünf Governance-Ebenen	25
8.4 Die zentrale Unterscheidung: Autonomy Level vs. Trust Level	26
8.5 Governance-Matrix	27
8.6 Agent Governance und Auditierung.....	28
9. Service Operating Contracts.....	29
9.1 Definition	29
9.2 Aufbau und Steuerungskette	29
9.3 Beispiel eines Contracts.....	30
9.4 Autonomie-Modus	31
9.5 Verwendung durch Agenten	32
10. AI Service Brain.....	33
10.1 Rolle und Vision.....	33
10.2 Architektur und Speichermodell	34
10.3 Wissensebenen und Bestandteile	34
10.4 Learning Loop und Learning Agent.....	35
10.5 Retrieval und Ranking — wie das Brain Wissen findet	36
10.6 Das Brain als eigenständige, generalisierbare Einheit	37
11. Skill Catalog	39
11.1 Definition und Motivation	39
11.2 Skill-Struktur	39
11.3 Skill-Kategorien.....	39
11.4 Reifegrad und Governance von Skills	39
12. Pattern Library.....	41
12.1 Definition.....	41
12.2 Aufbau eines Patterns	42
12.3 Pattern-Kategorien (Auszug)	43
13. Technologiestack	44
13.1 Auswahlkriterien	44
13.2 Überblick	44
13.3 Die Rolle von n8n.....	45
13.4 Bewusste Multi-Vendor-Entscheidung.....	45
14. Technischer Startplan.....	46

14.1 Grundprinzip.....	46
14.2 Phasenmodell	46
14.3 Erstes Datenmodell	48
14.4 Erfolgskriterien der Studie.....	48
15. Forschungs- und Messmodell.....	49
15.1 Forschungsansatz	49
15.2 Metrikmodell je Hypothese.....	51
15.3 Bewertung von Erkenntnissen	51
15.4 Nachvollziehbarkeit durch ADRs	51
15.5 Zwischenergebnisse des Demonstrators (Stand Juni 2026).....	52
15.6 Systematische Messreihen über alle fünf Hypothesen (Stand Juni 2026, F-011)	58
15.7 Findings im Überblick (F-001 – F-011).....	60
16. Autonomous Service Economics	65
16.1 Die ökonomische Grundthese	65
16.2 Kostenarten im Vergleich	65
16.3 Wirtschaftliche Steuerung im System	65
16.4 Der eigentliche Vermögenswert	65
17. Risiken, Grenzen und offene Fragen	66
17.1 Risiken und Gegenmaßnahmen	66
17.2 Grundsätzliche Grenzen	66
17.3 Offene Forschungsfragen (Zusammenfassung).....	66
18. Roadmap 2026–2030	68
18.1 Framework-Ausblick: das System selbst als AI-operable Plattform öffnen	68
19. Handlungsempfehlungen	70
20. Fazit	71
Konkrete nächste Schritte	71
21. Glossar	74
Versionshistorie.....	76
Quellen- und Dokumentenverzeichnis.....	81

Abbildungsverzeichnis

Abbildung 1: Verantwortungsteilung zwischen Mensch und KI-Agentensystem.	9
Abbildung 2: Die sechs Perspektiven der Studie als Mindmap.	12
Abbildung 3: Agentenarchitektur — Anatomie einer Agentenentscheidung unter Governance, eingebettet in das Agenten-Ensemble.....	14
Abbildung 4: Ebenenmodell der AI-Operable-IT-Services-Referenzarchitektur.	17
Abbildung 5: Referenzarchitektur des Demonstrator-Systems.	19
Abbildung 6: Die acht Dimensionen des AI-Readiness-Modells.	21
Abbildung 7: Governance-Ebenen und Delegierbarkeit von Strategie bis Execution.	25
Abbildung 8: Positionierung typischer Aktionen im Spannungsfeld von technischer Autonomie und organisatorischem Vertrauen. Aktionen unten rechts („riskante Lücke“) sind technisch möglich, aber organisatorisch (noch) nicht freigegeben.....	27
Abbildung 9: Steuerungskette und Aufbau eines Service Operating Contracts.	29
Abbildung 10: Schrittweise Steigerung des Autonomiegrades.	31
Abbildung 11: Architektur und hybrides Speichermodell des AI Service Brain.	34
Abbildung 12: Forschungs- und betriebsbezogener Lernfluss münden in Lessons Learned, Patterns, Skills und das Referenzmodell.	36
Abbildung 13: Zusammenspiel von Agent, Skill, Pattern, Service Operating Contract und AI Service Brain.	42
Abbildung 14: Forschungsmodell — von den Hypothesen H1–H5 über Experimente, Messungen, Evaluationen und Findings zu Lessons Learned, die Referenzmodell, Pattern Library, Skill Catalog, Governance- und AI-Readiness-Modell sowie das AI Service Brain fortschreiben.	50
Abbildung 15: Roadmap der Studie von 2026 bis 2030.	68

1. Executive Summary

Künstliche Intelligenz verändert die Art und Weise, wie Software entwickelt, betrieben und genutzt wird. Die meisten aktuellen Initiativen konzentrieren sich darauf, **einzelne Tätigkeiten** effizienter zu machen — etwa die Codeerstellung zu beschleunigen, Betriebsdaten zu analysieren oder Tickets zu klassifizieren. Diese Studie verfolgt einen grundlegend weitergehenden Ansatz: Sie untersucht nicht, wie KI einzelne Aufgaben unterstützt, sondern **wie IT-Anwendungen und IT-Services künftig gestaltet werden müssen, damit KI-Agentensysteme große Teile ihres Lebenszyklus eigenständig steuern können** — von Betrieb über Optimierung bis zur Weiterentwicklung.

Im Mittelpunkt steht der Begriff der **AI-Operable IT-Services**: Anwendungen und Services, die von Beginn an so konzipiert werden, dass sie durch KI-Agenten verstanden, überwacht, betrieben, optimiert und weiterentwickelt werden können.

Die entscheidende Erkenntnis dieser Studie ist, dass die Innovation nicht im einzelnen Agenten liegt. Ein leistungsfähiger Agent allein erzeugt weder Vertrauen noch nachhaltigen Wert. Wert entsteht erst durch das Zusammenspiel von vier Bausteinen:

- einem **AI-operablen Service**, der maschinenlesbar, observierbar und über klar definierte Schnittstellen steuerbar ist;
- einer **Governance**, die regelt, *was* ein Agent darf — getrennt von dem, was er technisch *kann*;
- einem **Agentensystem**, das spezialisierte Fähigkeiten orchestriert; und
- einem **AI Service Brain**, das Wissen, Muster und Erfahrungen dauerhaft speichert und wiederverwendbar macht.

Aus dieser Einsicht leitet die Studie ein konsistentes Rahmenwerk ab. Es umfasst ein **fünfschichtiges Referenzmodell**, ein **AI-Readiness-Modell mit acht Dimensionen** zur Bewertung der Eignung von Systemen, ein **Governance- und Trust-Modell**, das technische Autonomie (*Autonomy Level*) konsequent von organisatorischem Vertrauen (*Trust Level*) trennt, sowie **Service Operating Contracts** als maschinenlesbare Betriebsverträge. Wissen wird im **AI Service Brain** gespeichert und über einen **Skill Catalog** und eine **Pattern Library** wiederverwendbar gemacht. Eine kontinuierliche **Service Evolution** schließt den Kreis: Agenten betreiben Services nicht nur, sie entwickeln sie weiter.

Der vorgeschlagene Weg ist evolutionär, nicht disruptiv. Agenten beginnen im reinen Beobachtungsmodus (observe) und steigern ihre Autonomie schrittweise über analyze, propose, approval bis autonomous — stets innerhalb maschinenlesbarer Leitplanken und vollständig auditierbar. Die menschliche Gesamtverantwortung bleibt zu jedem Zeitpunkt erhalten; sie verlagert sich von der Durchführung operativer Tätigkeiten hin zu Strategie, Governance, Compliance und Risikomanagement.

Dieses Whitepaper konsolidiert die Konzeptdokumente der Machbarkeitsstudie zu einem zusammenhängenden strategischen Bild. Es richtet sich an IT-Management, Architektur, Softwareentwicklung und Betrieb und liefert sowohl das konzeptionelle Fundament als auch einen konkreten technischen Startplan und eine Roadmap bis 2030.

2. Ausgangssituation und Motivation

2.1 Der aktuelle Stand: KI als Werkzeug

In den meisten Organisationen wird Künstliche Intelligenz heute als **Werkzeug zur Unterstützung einzelner Tätigkeiten** eingesetzt. Entwickler nutzen Coding-Assistenten, Betriebsteams setzen ML-gestützte Anomalieerkennung ein, Servicedesks automatisieren die Ticketklassifizierung. Diese Anwendungen sind wertvoll, bleiben aber **punktuell**: Sie beschleunigen bestehende Prozesse, ohne deren Struktur grundlegend zu verändern. Der Mensch bleibt der zentrale Akteur, der jede Aufgabe initiiert, überwacht und abschließt.

2.2 Die strukturelle Grenze

Diese punktuelle Automatisierung stößt an eine strukturelle Grenze. Klassische IT-Systeme wurden **für menschliche Betreiber** entwickelt. Ihr Betriebswissen liegt in Wikis, Runbooks, Architekturdiagrammen und vor allem in den Köpfen einzelner Personen. Es ist für Menschen lesbar, für Maschinen jedoch nur schwer nutzbar. Solange dies so bleibt, kann ein KI-Agent ein System bestenfalls *beobachten* — er kann es nicht *verstehen* und schon gar nicht eigenständig *verantworten*.

Die Folge: Der wirtschaftliche Hebel der KI bleibt begrenzt. Operative Tätigkeiten — Monitoring, Incident-Bearbeitung, Patching, Kostenkontrolle, Recovery — binden weiterhin erhebliche menschliche Ressourcen. Die Weiterentwicklung von Systemen bleibt langsam und teuer.

2.3 Der Perspektivwechsel

Diese Studie vollzieht einen Perspektivwechsel. Sie fragt nicht: *Wie kann KI Menschen beim Betrieb helfen?*, sondern: *Wie müssen Systeme beschaffen sein, damit KI sie selbst betreiben kann?*

Damit rückt die **Gestaltung der Anwendung** selbst in den Mittelpunkt. Autonome IT-Services entstehen nicht allein durch bessere Agenten, sondern durch Anwendungen, die explizit für die Zusammenarbeit mit Agentensystemen entworfen werden. Diese Eigenschaft bezeichnen wir als **AI-Operability**.

Hypothese (H-AIR): Die Qualität einer IT-Anwendung wird künftig nicht nur daran gemessen, wie gut sie von Menschen betrieben werden kann, sondern auch daran, wie gut sie von KI-Agenten verstanden, gesteuert, überwacht und weiterentwickelt werden kann.

2.4 Warum jetzt

Drei Entwicklungen machen die Untersuchung gerade jetzt sinnvoll:

Treiber	Beobachtung	Konsequenz
Reasoning-fähige Modelle	Modelle der aktuellen Generation können mehrstufig planen, Werkzeuge aufrufen und Ergebnisse bewerten.	Agenten können komplexe Betriebs- und Entwicklungsaufgaben übernehmen, nicht nur Textgenerierung.
Cloudnative Standardisierung	API-First, Infrastructure as Code, GitOps und OpenTelemetry sind etablierte Praxis.	Die technischen Voraussetzungen für maschinensteuerbare Systeme sind weitgehend vorhanden.
Wirtschaftlicher Druck	Betriebskosten und Fachkräftemangel steigen, Time-	Der Bedarf an skalierbaren, weniger personalintensiven

	to-Market verkürzt sich.	Betriebsmodellen wächst.
--	--------------------------	--------------------------

Die Kombination dieser Treiber eröffnet ein Zeitfenster, in dem das Konzept AI-operabler IT-Services praktisch erprobt werden kann.

3. Vision und Zielbild

3.1 Vision

Die langfristige Vision der Studie ist die Entwicklung einer neuen Form der IT-Organisation, in der KI-Agentensysteme große Teile der operativen und analytischen Tätigkeiten übernehmen, während Menschen sich auf strategische, wirtschaftliche und regulatorische Entscheidungen konzentrieren.

Die Vision ist ausdrücklich nicht die Ablösung des Menschen. Sie ist die **Verlagerung menschlicher Arbeit** von operativen hin zu strategischen Tätigkeiten.

3.2 Verantwortungsteilung im Zielbild

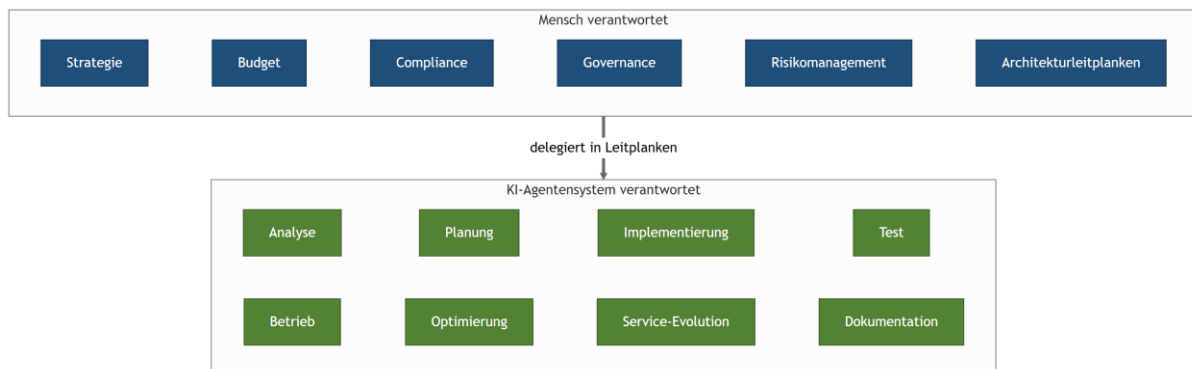


Abbildung 1: Verantwortungsteilung zwischen Mensch und KI-Agentensystem.

Die folgende Tabelle konkretisiert das Zielbild der Verantwortungsteilung. Die Prozentangaben beschreiben einen **angestrebten Zustand** auf hohem Reifegrad und sind als Hypothese zu verstehen, nicht als gemessenes Ergebnis.

Bereich	Mensch	KI-Agentensystem
Strategie	100 %	0 %
Budget	100 %	0 %
Compliance	100 %	0 %
Governance	90 %	10 %
Architekturprinzipien	80 %	20 %
Fachliche Anforderungen	70 %	30 %
Lösungsdesign	20 %	80 %
Implementierung	5 %	95 %
Test	5 %	95 %
Betrieb	1 %	99 %
Optimierung	10 %	90 %

Hypothese (H-ZIEL): Auf einem hohen AI-Readiness- und Trust-Reifegrad lässt sich der überwiegende Anteil von Betrieb, Implementierung und Test an Agentensysteme delegieren, ohne die menschliche Gesamtverantwortung aufzugeben.

3.3 Die drei parallel entstehenden Artefakte

Die Studie entwickelt bewusst **drei Artefakte parallel**. Keines davon ist für sich allein das Ziel — ihr Zusammenspiel ist es:

1. **Referenzsystem** — ein technischer Demonstrator (cloudnative KI-Anwendung mit Agentensystem).
2. **Referenzmodell** — die konzeptionelle Grundlage (Architektur, AI-Readiness, Governance).
3. **AI Service Brain** — das langfristige Wissens- und Fähigkeitsmodell.

Während Technologien, Modelle und konkrete Agenten über die Zeit ausgetauscht werden, bleiben Referenzmodell und AI Service Brain als dauerhafte Vermögenswerte bestehen.

Ende der Leseprobe

Dies ist ein Auszug (Deckblatt, Inhaltsverzeichnis, Kapitel 1-3) aus dem Whitepaper "AI-Operable IT-Services" von Markus Berg (Hessischer Rundfunk, Mitglied der ARD).

Die vollstaendige Studie (81 Seiten) koennen Sie per E-Mail anfordern:

`contact@agent-007.ai`

Mehr zur Praxisreferenz: <https://agent-007.ai/praxis/ai-operable-it-services/>