



HANDSON · FOUNDATIONAL WHITEPAPER

The HandsOn AI Operating Model

Sechs Design-Domänen für Organisationen, in denen KI wirklich funktioniert — ein konzeptionelles Framework, um KI durch Organisationsdesign wirksam zu machen.

PUBLIKATION

21. Juli 2026

VERSION

1.0

AUTOR

Maximilian Stein

HERAUSGEBER

HandsOn — AI
Strategy &
Governance
Consulting

KURZFASSUNG

KI scheitert ohne organisatorische Veränderung. Genau auf diese Veränderung sind wir spezialisiert.

Die meisten Organisationen investieren massiv in KI — doch laut einer aktuellen Deloitte-Studie skalieren weniger als 25 % erfolgreich über Piloten hinaus. KI-Initiativen bleiben den Nachweis eines nachhaltigen Effekts auf das Ergebnis weiterhin schuldig. Die bloße Einführung neuer Technologie reicht nicht aus, um echten Geschäftswert zu schaffen. Der Fokus muss sich auf die Gestaltung neuer Operating Models verlagern, die Hand in Hand mit den neuesten technologischen Entwicklungen arbeiten — und flexibel genug bleiben, um mit künftigen Entwicklungen Schritt zu halten.

Dieses Whitepaper stellt das HandsOn AI Operating Model vor: ein strukturiertes Framework aus sechs Design-Domänen, die gemeinsam definieren, wie sich eine Organisation weiterentwickeln muss, damit KI echten, nachhaltigen Geschäftswert liefert.

25%

haben >40 % ihrer KI-Piloten in Produktion gebracht.

16%

haben Jobs um KI-Fähigkeiten herum neu gestaltet.

21%

haben reife Governance-Modelle für KI-Agenten.

3,5×

mehr EBIT-Wirkung, wenn Workflows neu gestaltet werden.

QUELLEN Deloitte State of Generative AI 2026 · Deloitte AI Institute 2026 · McKinsey State of AI 2025.

EXECUTIVE SUMMARY · FÜNF BEFUNDE

01 Der Engpass ist nicht die Technologie — es ist die Organisation um sie herum.

Die meisten Unternehmen haben Jobs, Entscheidungsstrukturen und Operating Models nicht für KI neu gestaltet. Strategie, Struktur, Governance, Entscheidungen, Prozesse und Fähigkeiten müssen sich gemeinsam bewegen.

02 Die Pilot-Falle ist strukturell — kein Umsetzungsfehler.

Ein Pilot läuft mit einem kleinen Team und einem sauberen Datensatz. Produktion braucht Infrastruktur, Sicherheit, Compliance und Monitoring. Die meisten Organisationen bleiben stecken, weil ihnen das Operating Model für Skalierung fehlt.

03 Vier organisatorische Treiber erklären die Lücke.

Unklare Mensch-KI-Entscheidungsrechte · kein AI Ownership Model · KI auf bestehende Prozesse aufgesetzt · deutlich unterentwickelte KI-Fähigkeiten.

04 Das HandsOn AI Operating Model ist die strukturierte Antwort.

Sechs interdependente Design-Domänen auf zwei Ebenen, zusammengehalten von der Kernidee des Human-AI Interface als zentralem Element operativer Gestaltung im KI-Zeitalter.

05 Methodisch fundiert — Galbraith, Kates & Kesler, Worren, van Bree.

Bewährte Organisationsdesign-Methodik, erstmals in den KI-nativen Kontext erweitert — die Organisation bewusst an erster, die Technologie an zweiter Stelle.

INHALT

Was dieses Papier behandelt.

EINFÜHRUNG

- **Kurzfassung & Executive Summary** 02
Die Herausforderung · die Pilot-Falle · die organisatorische Lücke · die HandsOn-Antwort

TEIL I · DIE PRÄMISSE

- 01 **Das Problem, das niemand löst** 05
Vorwort und die grundlegende Prämisse
- 02 **Die Architektur — zwei Ebenen, sechs Domänen, ein Kern** 07
Warum Technologie ein Fundament ist, keine Domäne

TEIL II · DER KERN

- 03 **Das Human-AI Interface** 10
Vier Designfragen · vier Autonomiestufen

TEIL III · DER FOUNDATION LAYER

- 04 **D01 — Strategy & Value Architecture** 14
Drei Investment-Orientierungen · Cost of Autonomy
- 05 **D02 — Organisationsstruktur** 16
CoE · Föderal · Center-Led Hybrid
- 06 **D03 — System Governance** 17
Risiko-Tiering · AI Ownership · Lifecycle & Feedback-Loop

TEIL IV · DER ACTIVATION LAYER

- 07 **D04 — Decision Architecture** 19
Decision Rights Registry · Override · Rekalibrierung
- 08 **D05 — Process & Workflow Architecture** 20
Overlay · Integrated Redesign · AI-First
- 09 **D06 — Capabilities & Culture** 21
Fünf Zielgruppen · zwölf bis vierundzwanzig Monate

TEIL V · REIFE & INTEGRATION

- 10 **Das Reife-Kontinuum** 23
Stufen 0–3 · Multidimensionales Reifeprofil · Transition Architecture
- 11 **Die Interdependenz der Domänen** 26
Warum Integration entscheidend ist

SCHLUSS

- 12 **Der organisatorische Imperativ** 27
Über HandsOn · Zusammenarbeit · der Autor



Der Engpass war nie die Technologie.

Jedes Gespräch über KI-Transformation kommt irgendwann an derselben unbequemen Beobachtung an: Die Organisationen, die am meisten in KI-Technologie investieren, sind nicht die, die am meisten davon profitieren.

ABSCHNITT 01 · VORWORT

Das Problem, das niemand löst.

Warum ein Markt voller KI-Strategien, Governance-Frameworks und Reifegradmodelle die zentrale Frage unbeantwortet lässt.

Jedes Gespräch über KI-Transformation kommt irgendwann an derselben unbequemen Beobachtung an: Die Organisationen, die am meisten in KI-Technologie investieren, sind nicht die, die am meisten davon profitieren. Der Engpass ist selten die Technologie. Er ist fast nie die Strategie. Er ist mit bemerkenswerter Konsistenz die Organisation um die Technologie herum — die Struktur, die nie für KI neu gestaltet wurde; die Governance, die nachträglich angebaut wurde; die Entscheidungsarchitektur, die noch immer annimmt, dass jede Entscheidung durch einen Menschen läuft; die Belegschaft, die auf ein Tool geschult, aber nie darauf vorbereitet wurde, mit einem intelligenten Agenten zusammenzuarbeiten.

Der Beratungsmarkt hat auf dieses Problem mit einer Flut von KI-Strategien, KI-Governance-Frameworks, KI-Reifegradmodellen und KI-Readiness-Assessments reagiert. Jedes davon ist nützlich. Keines davon ist ausreichend. Sie adressieren Symptome statt Ursachen. Sie helfen Organisationen zu verstehen, wo sie stehen — aber nicht, wie sie neu gestalten, was sie sind. Und sie tun das, ohne auf den einen Wissensbestand zurückzugreifen, der für das Gestalten von Organisationen tatsächlich existiert:

Organisationsdesign als rigorose, erprobte, methodisch fundierte Disziplin.

DIE HANDSON-POSITION

Das HandsOn AI Operating Model ist unsere Antwort auf diese Lücke. **Es ist kein weiteres Reifegradmodell. Es ist keine Governance-Checkliste.** Es ist ein strukturiertes, operativ spezifisches Framework, um eine Organisation so umzugestalten, dass KI in ihr tatsächlich funktionieren kann — skaliert, nachhaltig und ohne den Zyklus aus Piloten, die nie Produktion werden, und Investitionen, die nie Wert werden.

DIE GRUNDLEGENDE PRÄMISSE

KI-Transformation ist ein Organisationsdesign-Problem, kein Technologie-Rollout-Problem. Diese Aussage klingt intuitiv, wird in der Praxis aber von fast jeder großen KI-Initiative widerlegt, die wir beobachtet haben. Organisationen behandeln KI konsequent als Projekt, das gemanagt wird, als Tool, das ausgerollt wird, als Fähigkeit, die eingekauft wird. Sie investieren in Modelle, Plattformen, Infrastruktur und Talente — und stellen dann fest, dass nichts davon Wert auf Unternehmensebene erzeugt, weil das organisatorische System darum herum nie dafür gestaltet wurde.

„Die Leistungslücke zwischen Organisationen, die KI als Werkzeug zu einer bestehenden Organisation hinzufügen, und jenen, die die Organisation um das Potenzial von KI herum neu gestalten, ist messbar, wächst — und wird sich ohne bewusste Intervention nicht schließen.“

— DIE GRUNDLEGENDE PRÄMISSE

ABSCHNITT 01 · FORTSETZUNG

In Organisationsdesign verankert — in den KI-nativen Kontext erweitert.

Die Evidenz ist eindeutig. Laut Deloittes State-of-AI-Forschung 2026 haben weniger als 25 % der Organisationen 40 % oder mehr ihrer KI-Experimente erfolgreich in Produktion gebracht. **84 %** haben Jobs nicht um KI-Fähigkeiten herum neu gestaltet. Nur **21 %** haben ein reifes Governance-Modell für autonome Agenten. McKinseys Daten von 2025 zeigen: Organisationen, die mehr als 5 % ihres EBIT auf KI zurückführen — die echten High Performer —, haben ihre Workflows mit **dreieinhalbfacher** Wahrscheinlichkeit fundamental neu gestaltet, statt bestehende nur zu ergänzen.

Das HandsOn AI Operating Model existiert, um genau dieses Redesign zu ermöglichen. Es gründet auf der Methodik von Jay Galbraiths Star Model, den Center-Led-Operating-Model-Strukturen von Amy Kates und Greg Kesler, den axiomatischen Designprinzipien von Nicolay Worren und der Organisationsdesign-Arbeit von Jeroen van Bree — rigoros adaptiert und für den KI-nativen Kontext erweitert. Nach unserem Wissen greift kein bestehendes Beratungs-Framework in dieser Weise auf Organisationsdesign-Methodik zurück.

Das ist keine theoretische Unterscheidung. Es bedeutet, dass jede Designentscheidung, die das Framework empfiehlt, auf eine zugrunde liegende Logik darüber zurückführbar ist, wie organisatorische Systeme tatsächlich funktionieren — nicht bloß auf das, was in diesem konkreten technologischen Moment zufällig funktioniert.

METHODISCHE FUNDAMENTE

Jay Galbraith — das Star Model aus fünf interdependenten Designhebeln.

Amy Kates & Greg Kesler — Center-Led-Operating-Model-Strukturen.

Nicolay Worren — axiomatische Designprinzipien für Organisationen.

Jeroen van Bree — Organisationsdesign für Capability Building.

„84 % — 21 % — 3,5×: haben Jobs nicht neu gestaltet · nur 21 % haben reife Governance für autonome Agenten · Top-Performer haben mit 3,5-facher Wahrscheinlichkeit ihre Workflows neu gestaltet.“

DELOITTE AI INSTITUTE 2026 · MCKINSEY STATE OF AI 2025

ABSCHNITT 02 · ABBILDUNG 01

Zwei Ebenen. Sechs Domänen. Ein Kern.

Jedes Element ist notwendig. Keines ist für sich allein ausreichend. Das Framework ist explizit holistisch: Eine Organisation, die fünf der sechs Domänen adressiert und die sechste vernachlässigt, erreicht nicht die Integration, die die anderen fünf voraussetzen.

FOUNDATION LAYER · WAS DIE FÜHRUNG ARCHITEKTONISCH FESTLEGT



DER KERN · DESIGNOBJEKT

Human-AI Interface

Wer entscheidet · Wer verantwortet · Wie das System lernt · Wo die Grenzen liegen

L1 · HUMAN-IN-THE-LOOP
KI berät

L2 · SUPERVISED EXECUTION
KI entscheidet, Mensch prüft

L3 · MONITORED AUTONOMY
KI entscheidet, Mensch wird informiert

L4 · HUMAN-IN-THE-EXCEPTION
Agent handelt; Mensch bei Anomalie

ACTIVATION LAYER · WAS MENSCHEN ERLEBEN



METHODIK · GALBRAITH STAR · KATES & KESLER · WORREN · VAN BREE

ABB. 01 Das HandsOn AI Operating Model — sechs Design-Domänen, zusammengehalten vom Human-AI Interface im Kern.

TECHNOLOGIE ALS FUNDAMENT, NICHT ALS DOMÄNE

Das HandsOn AI Operating Model enthält keine Technologie-Domäne. Technologieentscheidungen sind Voraussetzungen, die Organisationsdesign ermöglichen — sie sind nicht das Design selbst. Technologie entwickelt sich schneller als jede andere Dimension. Die HandsOn-Position: **Technologie ist eine notwendige Ermöglichungsbedingung, keine Design-Domäne.**

ABSCHNITT 02 · FORTSETZUNG

Foundation, Activation, Interface.

Was die Führung von oben architektonisch festlegt. Was Menschen in ihrer täglichen Arbeit erleben. Und die Grenze, die beides verbindet.

Die zwei Ebenen markieren einen fundamentalen Unterschied in der Natur der Designarbeit. Der **Foundation Layer** adressiert, was die Führung architektonisch festlegen muss — die Entscheidungen, die das gesamte organisatorische System von oben prägen. Der **Activation Layer** adressiert, was Menschen tatsächlich erleben — die operative Realität, jeden Tag in einer KI-fähigen Organisation zu arbeiten.

Beide Ebenen sind unverzichtbar. Organisationen, die nur in den Foundation Layer investieren, produzieren Governance-Frameworks und Strategien, die in der Dokumentation existieren, aber nicht in der Praxis. Organisationen, die nur in den Activation Layer investieren, erzeugen operative Veränderungen, die sich nicht selbst tragen können — weil ihnen die strukturelle und strategische Kohärenz zum Skalieren fehlt.

Im Zentrum beider Ebenen, mit jeder Domäne verbunden, sitzt das **Human-AI Interface** — das zentrale Designobjekt des Frameworks und das Konzept, das den HandsOn-Ansatz am deutlichsten von allem unterscheidet, was derzeit am Markt verfügbar ist.

WARUM KEINE TECHNOLOGIE-DOMÄNE

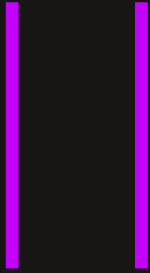
Die meisten KI-Frameworks enthalten eine Technologie-Ebene: Modellauswahl, Infrastruktur, Datenarchitektur, MLOps, Plattformsentscheidungen. Sie als Design-Domäne aufzunehmen, erzeugt zwei strukturelle Probleme.

Erstens folgen Technologieentscheidungen nicht derselben organisatorischen Designlogik wie Struktur-, Governance- oder Fähigkeitsentscheidungen. Eine Datenplattform ist kein organisatorisches Designobjekt im selben Sinne wie eine Governance-Architektur oder ein Entscheidungsrechte-Framework.

Zweitens entwickelt sich Technologie schneller als jede andere Dimension dieses Frameworks. Ein Modell, das konkrete Technologieentscheidungen festschreibt, veraltet schneller, als seine organisatorischen Erkenntnisse angewendet werden können.

DIE ZWEI FRAGEN

„Welche LLM-Plattform sollten wir nutzen?“ — eine Beschaffungsfrage. „Wie governen wir die Entscheidungen, die unsere KI-Systeme auf dieser Plattform treffen?“ — eine Organisationsdesign-Frage. Das HandsOn AI Operating Model ist gebaut, um die zweite zu beantworten.



Das Human-AI Interface.

Es ist kein Software-Interface. Es ist keine UX-Frage. Es ist die organisatorische Architektur von Entscheidungsfindung und Verantwortung in einer Welt, in der Menschen und KI-Systeme als aktive Agenten in denselben Workflows arbeiten.

ABSCHNITT 03

Ein Designobjekt erster Klasse.

Wo immer ein KI-System und ein Mensch gemeinsam Verantwortung für ein Ergebnis tragen, existiert ein Human-AI Interface — ob es gestaltet wurde oder nicht.

Wenn ein KI-System eine Kreditempfehlung erzeugt und ein Sachbearbeiter sie freigibt, existiert ein Human-AI Interface. Wenn eine Qualitätsprüfungs-KI Serienchargen freigibt und eine Schichtleitung eine Tageszusammenfassung erhält, existiert ein Human-AI Interface. Wenn ein autonomer Beschaffungsagent innerhalb definierter Parameter Lieferantenbestellungen auslöst und ein Mensch nur bei Überschreiten eines Schwellenwerts benachrichtigt wird, existiert ein Human-AI Interface.

Die entscheidende Einsicht: **Dieses Interface existiert unabhängig davon, ob es gestaltet wurde.** In den meisten Organisationen wurde es nicht gestaltet. Es hat sich angesammelt — durch Einzelentscheidungen darüber, welche Tools eingeführt, welche Freigaben verlangt, welche Benachrichtigungen erzeugt werden — zu einem inkonsistenten, ungoverneten Flickwerk, das niemandem gehört und das niemand vollständig versteht.

EINE ECHTE INNOVATION IM OD-DENKEN

Über die gesamte Existenz des Organisationsdesigns als Disziplin — von Galbraiths frühen Arbeiten in den 1970ern bis zu den jüngsten theoretischen Entwicklungen — beschäftigte sich das Feld damit, wie Menschen innerhalb organisatorischer Strukturen zueinander in Beziehung stehen. Das Human-AI Interface führt erstmals ein Designobjekt ein, das beschreibt, wie Menschen und nicht-menschliche Agenten Autorität, Verantwortung und Lernen innerhalb desselben organisatorischen Systems teilen.

VIER DESIGNFRAGEN DEFINIEREN DAS INTERFACE

F1

Wer entscheidet?

Auf welcher Autonomiestufe handelt die KI — und wo steht der Mensch?

F2

Wer verantwortet?

Auch wenn die KI gehandelt hat: Verantwortung verschwindet nicht.

F3

Wie lernt das System?

Aus wessen Urteil verbessert sich das System — und wie wird das Signal erfasst?

F4

In welchen Grenzen?

Autonomie ohne Grenzen ist kein Design — sondern Kontrollverzicht.

F1 · WER ENTSCHEIDET — UND AUF WELCHER AUTONOMIESTUFE

Das Autonomie-Spektrum ist nicht binär.

Es zerfällt nicht sauber in „der Mensch entscheidet“ und „die KI entscheidet“. Es verläuft über vier operativ unterscheidbare Stufen — jede eine andere Verteilung von Entscheidungsautorität und ein anderes Set organisatorischer Anforderungen.

AUTONOMIE-SPEKTRUM · WER ENTSCHEIDET & WIE

MENSCHLICHE AUTORITÄT

KI-AUTORITÄT

L1 · HUMAN-IN-THE-LOOP

KI berät. Der Mensch entscheidet immer.

Keine Wirkung ohne menschliche Bestätigung. Universeller Startpunkt. Die richtige Konfiguration für hochriskante, irreversible oder stark kontextabhängige Entscheidungen.

L2 · SUPERVISED EXECUTION

KI entscheidet. Der Mensch prüft innerhalb eines definierten Zeitfensters.

Der Default ist von menschlicher Bestätigung zu KI-Ausführung gewandert — eine qualitative Veränderung organisatorischer Verantwortung, die L1 nicht erfordert.

L3 · MONITORED AUTONOMY

KI handelt. Der Mensch wird informiert; greift bei Anomalien ein.

Menschliche Aufmerksamkeit wandert von der Entscheidungs- auf die Systemebene: nicht „Ist diese Entscheidung korrekt?“, sondern „Performt dieses System über seine gesamte Entscheidungspopulation?“

L4 · HUMAN-IN-THE-EXCEPTION

Autonome KI-Agenten führen ganze Workflows Ende-zu-Ende aus. Menschen definieren Ziele, Randbedingungen und Leistungsstandards — und werden nur bei Anomalien, Grenzverletzungen oder geplanten Reviews einbezogen. Das ist die Konfiguration, die Agentic AI erfordert — und sie verlangt die anspruchsvollste Governance-Architektur aller vier Stufen.

ABB. 02 Die vier Stufen zu definieren ist die konzeptionelle Arbeit des Human-AI Interface. Zu governen, welche konkreten Entscheidungen auf welcher Stufe operieren, ist die organisatorische Arbeit von D04 — Decision Architecture.

F2 · WER VERANTWORTET

KI absorbiert keine Verantwortung — sie verlagert nur die Versuchung, sie abzugeben.

Wenn ein KI-System einen Kreditantrag genehmigt, eine Produktionslinie stoppt oder eine Kundenbeschwerde weiterleitet, liegt die Versuchung nahe, die Verantwortung als an das System delegiert zu betrachten. **Das ist sie nicht.**

Verantwortung für KI-gestützte oder KI-generierte Entscheidungen bleibt bei Menschen — beim **AI Owner**, der das System beauftragt hat und governiert, beim **AI Steward**, der es technisch betreut, und in regulierten Domänen bei der Person, deren Name formal mit der Entscheidung verbunden ist.

Das Human-AI Interface macht diese Verantwortungskette explizit. Es benennt die Verantwortlichen pro System und pro Entscheidungstyp, definiert Eskalationspfade für strittige KI-Ausgaben — und stellt sicher, dass organisatorische Verantwortung nicht in der Annahme verschwindet, die Maschine sei zuständig.

F3 · WIE LERNT DAS SYSTEM — UND VON WEM

Der Feedback-Loop ist die am meisten unterschätzte Dimension.

Er entscheidet, ob sich die KI-Fähigkeit einer Organisation über die Zeit wirklich verbessert — oder statisch bleibt. KI-Systeme in Produktion sind nicht fix. Sie driften, wenn sich Datenverteilungen verändern. Sie degradieren, wenn sich Bedingungen über ihr Training hinaus entwickeln. Sie verbessern sich — aber nur, wenn Feedback aus menschlichem Urteil systematisch erfasst und zurückgespielt wird.

Jeder menschliche Override ist ein Lernsignal. Jeder markierte Grenzfall enthält Information über Randbedingungen, die die KI noch nicht gelernt hat. Jede Eskalation zeigt dem System, wo seine Konfidenzschwellen justiert werden müssen. Das Human-AI Interface gestaltet den Feedback-Loop.

D03 — System Governance benennt seinen organisatorischen Eigentümer: den AI Steward.

„Ohne diese Struktur werden Organisationen über die Zeit nicht KI-fähiger — sie betreiben schlicht dieselben Systeme in derselben Qualität auf alternden Annahmen, bis etwas schlimm genug schiefgeht, um ein Umdenken zu erzwingen.“

— ÜBER DEN FEEDBACK-LOOP

F4 · GRENZBEDINGUNGEN

Autonomie ohne Grenzen ist kein Design — sondern Kontrollverzicht.

Die vierte Dimension des Human-AI Interface betrifft die explizite Definition des operativen Rahmens, innerhalb dessen die Autorität eines KI-Systems gilt. Grenzbedingungen sind nicht dasselbe wie Governance-Regeln.

Governance-Regeln definieren, was ein KI-System grundsätzlich tun darf. Grenzbedingungen definieren die konkreten operativen Parameter — Datenumfang, Entscheidungsschwellen, Transaktionslimits, Eskalationstrigger, Zeitfenster —, innerhalb derer diese Autorisierung gilt. Fallen Bedingungen aus dem definierten Rahmen, muss das System eskalieren, anhalten oder an menschliches Urteil übergeben — unabhängig von seiner formalen Autonomiestufe.

Diese Frage wird umso folgenreicher, je höher die Autonomie. Auf **Stufe 1** sind Grenzbedingungen weitgehend implizit: Der Mensch, der jede Ausgabe prüft, bildet eine natürliche Eindämmungsschicht. Auf **Stufe 4**, wo autonome Agenten Ende-zu-Ende-Workflows ohne vorherige menschliche Freigabe ausführen, sind explizite Grenzbedingungen der primäre Mechanismus organisatorischer Kontrolle.

ZWEI FEHLERMODI AUF STUFE 4

Überängstliches Anhalten. Ein Agent stoppt unnötig bei Randfällen, für die er nie ausgestattet wurde.

Handeln außerhalb des Auftrags. Ein Agent agiert jenseits seines vorgesehenen Rahmens, weil kein explizites Limit gesetzt wurde.

Beides sind organisatorische Designfehler, keine Technologiefehler. Grenzbedingungen müssen so präzise gestaltet werden wie die Autonomiestufe selbst — und sie sind die primäre operative Verantwortung des AI Steward.

GRENZEN ALS FINANZIELLE GOVERNANCE-VARIABLE

Agenten mit unterspezifizierten Grenzen verbrauchen deutlich mehr Tokens, weil sie Unsicherheits- und Korrekturpfade durchlaufen. **Grenzdesign ist eine finanzielle Governance-Variable** — ein Punkt, der in D01 im Rahmen des Cost-of-Autonomy-Frameworks vertieft wird.

„Bevor irgendeine Domäne gestaltet werden kann, muss die Organisation ein gemeinsames Vokabular für das Verhältnis von menschlichem Urteil und KI-Handlungsfähigkeit etablieren.“

— WARUM DAS INTERFACE DER KERN IST

DOMÄNE 01 · FOUNDATION LAYER

Strategy & Value Architecture.

Wo schafft KI echten Wettbewerbsvorteil — und wo nur operative Effizienz? Eine Organisation, die beides nicht trennt, verwaltet ein inkohärentes Portfolio, statt eine kohärente Fähigkeit aufzubauen.

Eine Organisation, die KI primär zur Senkung von Prozesskosten in Backoffice-Funktionen nutzt, betreibt ein fundamental anderes KI-Programm — mit anderen organisatorischen Anforderungen — als eine, die mit KI eine Fähigkeit in prädiktiver Qualitätssicherung aufbaut, die Wettbewerber nicht replizieren können. Beides sind legitime strategische Entscheidungen. **Sie explizit zu treffen, statt sie implizit entstehen zu lassen, unterscheidet eine KI-Strategie von einem KI-Projektportfolio.**

DREI INVESTMENT-ORIENTIERUNGEN

Top-Down Strategic Mandate

Die Führung definiert zwei bis drei strategische KI-Wetten entlang der Wettbewerbspositionierung. Schafft Klarheit und Ownership — setzt aber hohe strategische Klarheit voraus und verstärkt strategische Fehler, wenn sie passieren.

Portfolio-Led Use-Case Management

Das Portfolio entsteht bottom-up: Use Cases werden systematisch identifiziert und nach Impact und Machbarkeit priorisiert. Erzeugt Beteiligung — führt ohne strategischen Rahmen aber zu Fragmentierung: viele Piloten, keine kohärente Fähigkeit.

Capability-Based Investment Thesis

EMPFOHLEN

Definiert drei bis fünf strategische KI-Fähigkeiten und behandelt Use Cases als Ausdruck dieser Fähigkeiten. Schafft die Bedingungen für wiederverwendbare KI-Assets, IP-Entwicklung und die organisatorische Kohärenz, die Vorteile nachhaltig macht.

IN DER PRAXIS

Die ersten beiden als Standard-Einstieg kombinieren: Top-down-Rahmung definiert das Terrain, plus Bottom-up-Portfoliomanagement. Die Capability-Based Investment Thesis wird zur dominanten Logik ab Reifestufe 2 — wenn der Produktions-Track-Record zeigt, welche Fähigkeiten wirklich differenzieren.

DELIVERABLE — DIE AI STRATEGY MAP

- Dokumentierte Portfolio-Klassifikation — strategisch vs. operativ.
- Eine Wertdefinition pro Initiative.
- Ein Build-/Buy-/Partner-Entscheidungsframework.
- Eine quartalsweise Governance-Kadenz für das Portfolio-Review.
- Ein **Cost-of-Autonomy-Profil** pro Initiative.

D01 · COST OF AUTONOMY

Strategische KI-Entscheidungen sind auch ökonomische Entscheidungen.

Die Kosten eines KI-Systems sind eine Funktion der Autonomiestufe, auf der es operiert — und die Kostenstruktur verschiebt sich fundamental, wenn die Autonomie steigt.

DOMINANTER KOSTENTREIBER · NACH AUTONOMIESTUFE

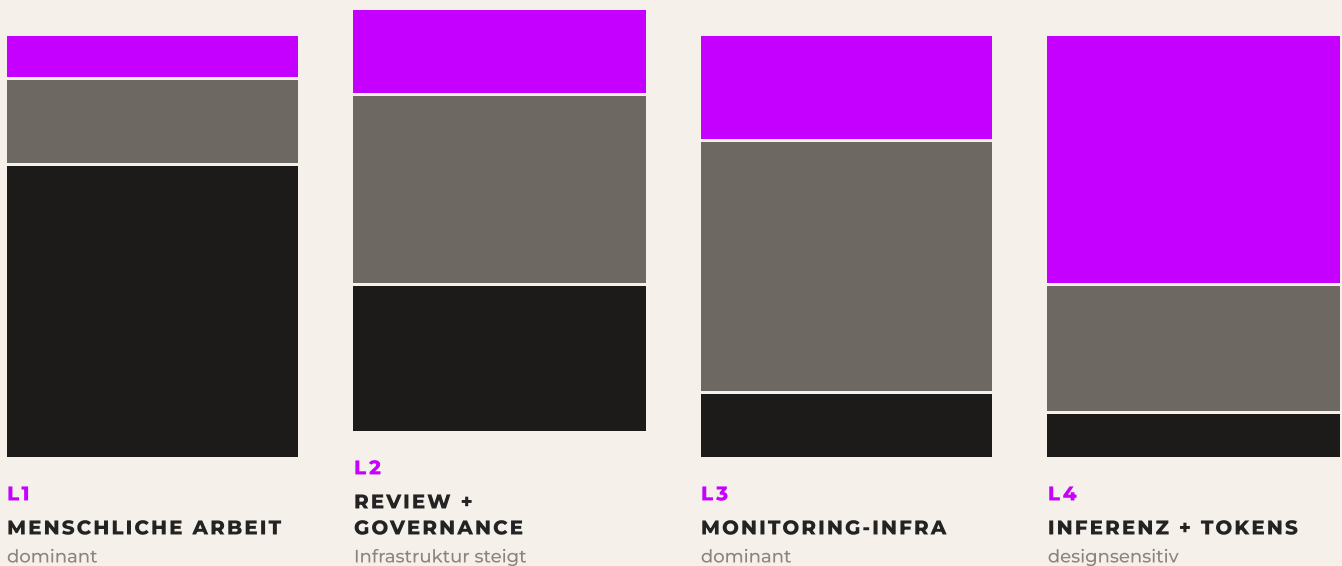


ABB. 03 Mit steigender Autonomie verändert sich die Kostenstruktur — nicht bloß das Kostenniveau. Auf L4 ist Designqualität eine finanzielle Variable.

Auf **Stufe 1** dominiert der Kostenfaktor menschliche Arbeit: Jede KI-Ausgabe durchläuft einen menschlichen Prüfzyklus. Auf den **Stufen 2 und 3** sinken die Arbeitskosten, aber die Kosten der Governance-Infrastruktur steigen — Monitoring-Systeme, Anomalieerkennung, Override-Logging und Rekalibrierungsprozesse.

Auf **Stufe 4** werden KI-Betriebskosten — Inferenzvolumen, mehrstufiges Reasoning, agentische Ausführung — dominant. Token-Kosten sind kein Fixposten, sondern eine dynamische operative Variable, die direkt auf Designqualität reagiert. **Die Qualität des Organisationsdesigns hat einen direkten, messbaren Effekt auf die KI-Betriebskosten.**

DOMÄNE 02 · FOUNDATION LAYER

Organisationsstruktur.

Die umstrittenste Designentscheidung des Frameworks — und die, bei der die meisten Organisationen ihre folgenreichsten Fehler machen.

ZENTRALES AI CENTER OF EXCELLENCE

Bündelt alle KI-Kompetenzen — Data Science, ML Engineering, KI-Produktmanagement, KI-Governance — in einer zentralen Einheit, die die Geschäftsbereiche bedient. Schafft einheitliche Standards, effiziente Ressourcennutzung und klare Governance-Verantwortung. **Funktioniert gut in frühen Phasen**, wenn die zentrale Herausforderung darin besteht, organisatorische Fähigkeit aufzubauen und Qualitätsstandards zu etablieren.

Mit zunehmender Reife wird es kontraproduktiv: Das CoE wird zur Warteschlange, ihm fehlt das Domänenwissen der Geschäftsbereiche, und es verhindert, dass lokale Ownership entsteht.

VOLLSTÄNDIG FÖDERALES MODELL

Verankert KI-Fähigkeit vollständig in den Geschäftsbereichen. Maximiert Domänenwissen und parallele Innovation, die ein zentrales Modell strukturell verhindert. In Organisationen ohne starke Governance-Kultur erzeugt es aber redundante Arbeit, inkonsistente Qualität und die Governance-Lücken, in denen Schatten-KI wuchert.

DER CENTER-LED HYBRID **HANDSON-ZIEL**

Ein kleiner strategischer AI Hub — typischerweise fünf bis fünfzehn Personen — kombiniert mit Embedded AI Leads in jedem größeren Geschäftsbereich. Der Hub setzt Standards, entwickelt Plattformen, verantwortet Governance und baut Playbooks. **Er befähigt, statt zu kontrollieren.** Embedded AI Leads treiben die Umsetzung, übersetzen Hub-Standards in den Domänenkontext und bilden die Brücke zwischen strategischer Kohärenz und operativer Geschwindigkeit.

DIE DOPPELTE BERICHTSLINIE

Jeder Embedded AI Lead berichtet **funktional** an die Bereichsleitung und **methodisch** an den AI Hub. Ohne diese doppelte Verantwortlichkeit kollabiert das Modell — entweder in ein Schatten-CoE (der Hub übernimmt Kontrolle) oder in ein föderales Modell mit nomineller Koordination (der Hub kann Standards nicht durchsetzen).

Struktur muss der Autonomie folgen. Eine Organisation auf L1 kann mit einem zentralen CoE funktionieren. Eine Organisation auf L3 über Kernprozesse hinweg kann es nicht — die Komplexität Dutzender Produktionssysteme übersteigt, was eine zentrale Einheit leisten kann, ohne zum Engpass zu werden.

DOMÄNE 03 · FOUNDATION LAYER

System Governance.

In zwei entgegengesetzte Richtungen gleichzeitig missverstanden: als Compliance-Checkbox oder als Innovationsbremse. Beide Deutungen erzeugen Architekturen, die scheitern.

Die korrekte Rahmung: **Governance ist ein Wachstumskatalysator**. Eine Organisation mit reifer Governance-Architektur kann KI schneller ausrollen als eine ohne — weil sie die Verantwortungsfragen, die vor jedem Deployment geklärt sein müssen, im Voraus geklärt hat; weil sie Risikobewertungsprozesse geschaffen hat, die risikoarme Initiativen schnell passieren lassen und risikoreiche gründlich prüfen; und weil sie das Vertrauen bei Legal, Compliance, Risk und den Mitarbeitenden aufgebaut hat, das KI-Initiativen Rückhalt statt Widerstand verschafft.

RISIKO-TIERING — DAS EU-AI-ACT-FUNDAMENT

Die Klassifikation jeder KI-Anwendung nach Risikostufe ist das Fundament. Der EU AI Act liefert dafür einen gesetzlich verankerten Rahmen — **Verboden** (Social Scoring, biometrische Echtzeitüberwachung), **Hochrisiko** (Kreditentscheidungen, Personalauswahl, kritische Infrastruktur), **Begrenztes Risiko** (Deepfakes, Chatbots, Empfehlungssysteme), **Minimales Risiko** (Produktivitäts-KI, Spamfilter, Dokumentenerstellung). Zentral ist das Verhältnismäßigkeitsprinzip: Risikoarme Systeme sollten minimalen Overhead tragen; Hochrisikosysteme brauchen umfassende Prüfung, Monitoring und Audit.

AI OWNERSHIP

Die Struktur, die Governance real statt theoretisch macht. Jedes produktive KI-System braucht einen benannten **AI Owner** — eine Führungskraft aus dem Business, verantwortlich für Outputs und die Entscheidungen, die sie speisen — und einen benannten **AI Steward** — eine technische Führungskraft, verantwortlich für Performance, Wartung und Verbesserung. Zusammen bilden sie eine Verantwortungskette vom Deployment über den Betrieb bis zur späteren Außerbetriebnahme.

LIFECYCLE GOVERNANCE & DER FEEDBACK-LOOP

KI-Systeme sind nicht statisch. Lifecycle Governance erstreckt Ownership über die gesamte Lebensdauer: laufendes Monitoring, periodische Risiko-Reklassifikation, systematische Erfassung von Override- und Eskalationsdaten sowie ein formaler Prozess für die Außerbetriebnahme. Der **AI Steward** ist der benannte Eigentümer des Feedback-Aggregationsprozesses — Signalerfassung, Aggregation und Review, Weiterleitung.

EU AI ACT · VIER RISIKOSTUFEN

STUFE 01 · VERBODEN

Untersagte Anwendungsfälle

Social Scoring · biometrische Echtzeitüberwachung im öffentlichen Raum · manipulative KI

STUFE 02 · HOCHRISIKO

Konformitätsbewertung

Kreditentscheidungen · Personalauswahl · kritische Infrastruktur · Medizin-KI · Strafverfolgung

Risikomanagement · Logs · menschliche Aufsicht · CE-Registrierung

STUFE 03 · BEGRENZTES RISIKO

Transparenzpflichten

Chatbots · Deepfakes · Empfehlungssysteme

Nutzer informieren · synthetische Inhalte kennzeichnen

STUFE 04 · MINIMALES RISIKO

Freiwillige Kodizes

Produktivitäts-KI · Spamfilter · Dokumentenerstellung

Keine Pflichtkontrollen — interne Governance gilt weiter

VERHÄLTNISSMÄSSIGKEIT · KONTROLLEN SKALIEREN MIT DEM RISIKO, NICHT MIT DER AMBITION

ABB. 04 EU-AI-Act-Risiko-Tiering. Jede Stufe trägt definierte Kontrollen und Gate-Prozesse, die ein System vor dem Deployment passieren muss.

IV

Was Menschen an jedem Arbeitstag erleben.

D04 Decision Architecture · D05 Process & Workflow Architecture · D06 Capabilities & Culture. Die drei Domänen, die die Architektur des Foundation Layer in operative Realität übersetzen.

DOMÄNE 04 · ACTIVATION LAYER

Decision Architecture.

Die am meisten vernachlässigte Domäne in der Arbeit an KI Operating Models — und die, an der die meisten Organisationen in der Praxis tatsächlich scheitern.

Der Unterschied zwischen D03 und D04 ist kategorisch, keine Frage des Grades. **D03 — System Governance** governt das KI-System als Objekt über seinen Lifecycle. **D04 — Decision Architecture** governt, was das System innerhalb dieses Lifecycles *zu entscheiden befugt ist*: welche Entscheidungstypen es ausführen darf, auf welcher Autonomiestufe, unter welchen Bedingungen, durch welchen Governance-Prozess. Eine Organisation kann reife System Governance und komplette Unordnung in ihrer Decision Architecture haben — und umgekehrt. Beides ist notwendig; keines ersetzt das andere.

DIE DECISION RIGHTS REGISTRY

Das zentrale Deliverable. Anders als ein einfaches Klassifikationsdokument ist die Registry ein **Autoritätsregister**: Für jeden bedeutsamen Entscheidungstyp legt sie fest, welche Funktion, welche Rolle und welche Governance-Stufe die Autorität hält, diesen Entscheidungstyp auf jeder Autonomiestufe zu betreiben. Sie dokumentiert nicht nur, wo eine Entscheidung liegt — sondern wer das entschieden hat, unter welchen Bedingungen, durch wen und auf welcher Evidenz.

Die Autonomiestufe ist das Ergebnis einer Autoritätsentscheidung, nicht bloß die Anwendung eines Frameworks. **Diese Unterscheidung gibt der Registry ihre organisatorische Durchsetzungskraft.**

KLASSIFIKATIONS-GOVERNANCE · WER ENTSCHEIDET, WIE ENTSCIEDEN WIRD

Der Aufbau der Registry erfordert einen formalen Prozess, der Entscheidungstypen an vier Fragen misst:

- Was steht bei einer falschen KI-Entscheidung auf dem Spiel?
- Ist die Entscheidung reversibel?
- Welche Qualität hat das KI-System nachweislich in Produktion?
- Ist die Governance-Infrastruktur bereit, die Aufsicht zu tragen?

OVERRIDE- & ESKALATIONSARCHITEKTUR

Die Echtzeit-Governance-Ebene. Für jede Autonomiestufe definiert D04, wer Override-Autorität hält, unter welchen Bedingungen sie ausgeübt werden *muss*, welches Logging erforderlich ist — und welche Verantwortung an Overrides hängt, die unterlassen wurden, obwohl sie geboten waren.

REKALIBRIERUNGS-GOVERNANCE

KI-Systeme sind nicht statisch. D04 definiert die strukturierte Kadenz — typischerweise **quartalsweise in den ersten zwei Jahren** — mit definierten Evidenzstandards für das Hochstufen eines Entscheidungstyps und definierter Autorität für das Herabstufen, wenn die Performance nachlässt.

DOMÄNE 05 · ACTIVATION LAYER

Process & Workflow Architecture.

Die größten Effizienzgewinne aus KI entstehen nicht, indem KI-Tools auf bestehende Prozesse gelegt werden. Sie entstehen, indem die Prozesse selbst um die spezifischen Fähigkeiten und Grenzen von KI herum neu gestaltet werden.

DREI PROZESSANSÄTZE · PARALLEL, NICHT SEQUENZIELL

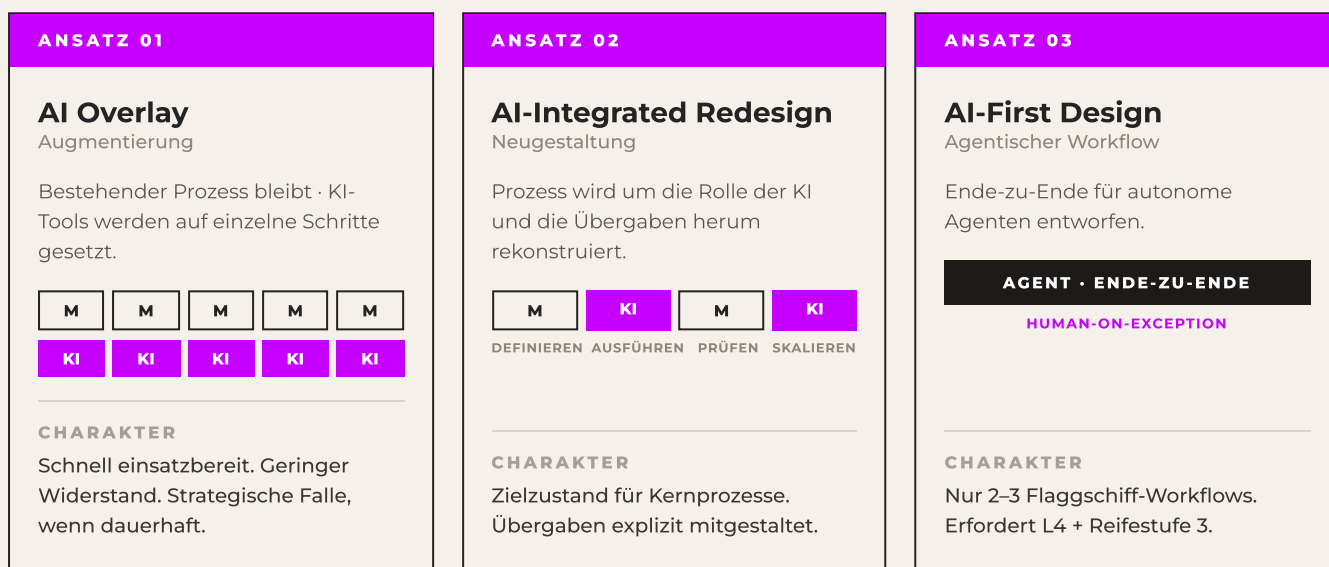


ABB. 05 Das HandsOn-Framework unterscheidet drei Prozessansätze, die in jeder reifen KI-fähigen Organisation parallel existieren.

AI Overlay ist schnell, widerstandsarm, der richtige Startpunkt für periphere Prozesse — und eine strategische Falle, wenn es als Dauerzustand behandelt wird. **AI-Integrated Redesign** ist der Zielzustand für Kernprozesse: Übergabepunkte als Designelemente erster Klasse, Qualitätssicherung und Ausnahmebehandlung eingebaut.

AI-First Design eignet sich für zwei bis drei Flaggschiff-Workflows, bei denen die Organisation die Governance, Datenreife und Reife besitzt, um auf L4 zu operieren. Die kritische Disziplin über alle drei Ansätze hinweg: **das explizite Design der Mensch-KI-Übergabepunkte** — dort passieren die meisten Fehler.

DOMÄNE 06 · ACTIVATION LAYER

Capabilities & Culture.

Die dauerhafteste Beschränkung der KI-Transformation ist nicht die Verfügbarkeit von Technologie, die Qualität der Dateninfrastruktur oder selbst Governance — sondern die Fähigkeit der menschlichen Organisation, in einem KI-fähigen Kontext wirksam zu arbeiten.

FÄHIGKEITS-ARCHETYP · NACH AUTONOMIESTUFE



FÜNF ZIELGRUPPEN · C-SUITE · MITTLERES MANAGEMENT · DOMÄNEN-EXPERTEN · TECHNISCHE KI-TEAMS · ALLE MITARBEITENDEN

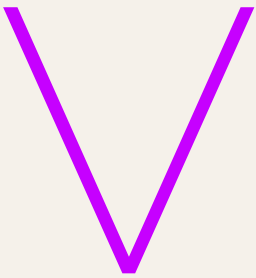
ABB. 06 Fähigkeits-Archetyp nach Autonomiestufe. Die L4-Fähigkeit ist eine völlig neue professionelle Kompetenz, die es vor der Ära agentischer KI in keiner nennenswerten Form gab.

McKinsey-Daten zeigen: Die Nachfrage nach KI-Kompetenz wuchs in zwei Jahren um das **Siebenfache**. Deloitte stellt fest, dass **84 %** der Organisationen Jobs nicht um KI-Fähigkeiten herum neu gestaltet haben. Die Lücke betrifft nicht primär technische Skills — es geht um die organisatorischen und verhaltensbezogenen Fähigkeiten, die auf L1–L4 erforderlich sind. Die HandsOn-Fähigkeitsarchitektur unterscheidet **fünf Zielgruppen** — C-Suite, mittleres Management, Domänen-Experten, technische KI-Teams, alle Mitarbeitenden — mit differenzierten Lernpfaden.

KULTURELLE TRANSFORMATION — ZWÖLF BIS VIERUNDZWANZIG MONATE

Vier kulturelle Marker von Organisationen, in denen KI-Transformation gelingt:

- Führungskräfte nutzen KI aktiv selbst;
- psychologische Sicherheit, Bedenken zu KI-Ausgaben zu äußern; Fehler werden als Lerngelegenheiten behandelt; ein kulturelles Narrativ von menschlicher Handlungsfähigkeit und Fähigkeitserweiterung — nicht Ersatz.



Reife ist eine Form, keine Zahl.

Organisationen stehen selten in allen sechs Domänen auf derselben Stufe. Das Multidimensionale Reifeprofil macht Ungleichgewichte auf Domänenebene sichtbar — und zeigt die gefährliche Lücke, wo reife Governance über Stufe-0-Prozessdesign sitzt.

ABSCHNITT 10

Vier Stufen — angewandt auf Domänenebene, nicht nur auf die Organisation.

Eine Organisation wird selten — wenn überhaupt — in allen sechs Domänen gleichzeitig auf derselben Stufe stehen. Die diagnostische Kraft liegt darin, diese Unterschiede auf Domänenebene sichtbar zu machen.

DIE VIER REIFESTUFEN



ABB. 07 Die vier Reifestufen. Nachhaltiger Betrieb auf jeder Autonomiestufe erfordert die entsprechende Reife-Infrastruktur.

Stufe 0 — Organische Emergenz ist die Stufe, die die meisten Frameworks auslassen und die die meisten Organisationen gerade durchleben. Mitarbeitende nutzen KI ohne Autorisierung oder Governance. Der primäre Fehlermodus: sie gar nicht als Stufe zu erkennen.

Stufe 1 — Augmentiert ist geprägt von beratenden KI-Beziehungen über alle sechs Domänen. Der primäre Fehlermodus ist **Pilotmüdigkeit**: viele kleine Initiativen, die lokalen Wert zeigen, aber nicht skalieren können.

Stufe 2 — Eingebettet: KI als struktureller Bestandteil von Kernworkflows. Center-Led Hybrid live, AI Ownership vergeben, Registry rekali­briert quartalsweise. Die primäre Herausforderung ist **Skalierung** — die Replikation über alle Einheiten.

Stufe 3 — Agentisch: Autonome KI-Agenten führen Ende-zu-Ende-Workflows aus, Menschen agieren als Orchestratoren und Ausnahme-Manager. Das häufigste Scheitern: **Stufe 3 ohne Stufe-2-Fundamente zu versuchen**.

ABSCHNITT 10 · FORTSETZUNG

Das ganzheitliche Reifeprofil ist eine vollständige Diagnose.

Eine als „Stufe 2“ bewertete Organisation kann Stufe-2-Governance und Stufe-0-Prozessdesign haben — eine strukturell gefährliche Konfiguration. Das aggregierte Label verbirgt das vollständig. Das Profil macht es sichtbar.

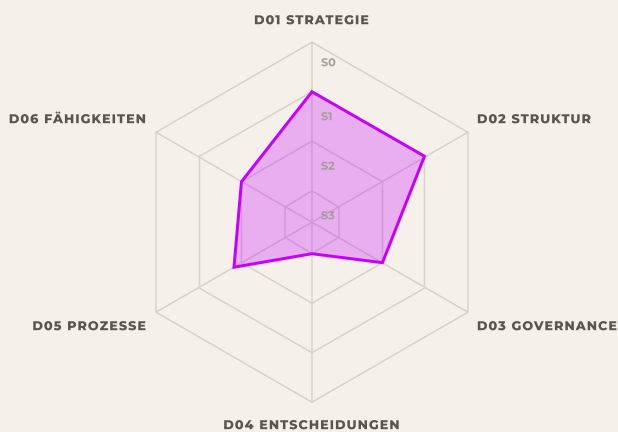


ABB. 08 Das Multidimensionale Reifeprofil wendet das Vier-Stufen-Kontinuum unabhängig auf jede der sechs Domänen an. Das Ergebnis ist keine Zahl — es ist eine Form.

DREI DIAGNOSTISCHE DOMÄNENPAARE

D03 & D04 — System Governance und Decision Architecture müssen parallel voranschreiten. Governance ohne Entscheidungsautorität produziert Compliance-Theater. Entscheidungsautorität ohne Governance produziert ungovernete Autonomie.

D01 & D05 — Strategie und Prozess müssen ausgerichtet sein. Eine ambitionierte AI Strategy Map über unveränderten Legacy-Prozessen ist keine Transformation — sie ist ein Wunschzettel.

D05 & D06 — Prozess-Redesign und Fähigkeitsentwicklung müssen sich gemeinsam bewegen. Neu gestaltete Prozesse funktionieren nur, wenn die Menschen an den Übergaben die Kompetenz haben, dort Urteilskraft auszuüben.

WARUM EINE FORM, KEINE ZAHL

Ein Profil, das Stufe 2 auf D03 und Stufe 0 auf D04 zeigt, ist kein nuanciertes Ergebnis, das Experteninterpretation braucht. **Es ist ein unmittelbares visuelles Signal**, dass die Organisation in die Governance ihrer KI-Systeme investiert hat, ohne je zu formalisieren, was diese Systeme zu entscheiden befugt sind. Diese Lücke ist handlungsleitend — und das Profil macht sie unübersehbar.

ABSCHNITT 10 · TRANSITION ARCHITECTURE

Die Reifestufen beschreiben Ziele. Die Transition Architecture beschreibt den Weg.

Die meisten KI-Transformationsprogramme scheitern nicht, weil der Zielzustand falsch wäre — sondern weil der Pfad zwischen den Stufen nie gestaltet wurde.

S0 → S1

Organisch zu strukturiert

Vorhandene informelle Fähigkeit sichtbar und governbar machen. Ein benannter KI-Sponsor auf Führungsebene. Eine minimal tragfähige Governance-Struktur. Ein bewusster Einsatz, die Arbeitspraktiken zu heben und zu teilen, die Early Adopter individuell entwickelt haben.

Fehlermodus: Überbauen —

umfassende Frameworks vor grundlegender Sichtbarkeit treiben die Nutzung weiter in den Untergrund.

S1 → S2

Augmentiert zu eingebettet

Drei nicht verhandelbare Enabler. **Strukturelle Evolution** zum Center-Led Hybrid mit doppelten Berichtslinien. **Prozess-Redesign** — mindestens ein Kernprozess wirklich rekonstruiert, nicht augmentiert. Die **Decision Rights Registry** reift von der Absichtserklärung zum operativen Instrument.

Primäre Herausforderung: eine Führungs-, keine Technologieherausforderung. Hier schließt sich die Pilot-Falle.

S2 → S3

Eingebettet zu agentisch

Ein qualitativer Sprung, keine höhere Intensität der S2-Arbeit. Vom Bestätigen von KI-Ausgaben zum Definieren von Zielen, Überwachen der System-Performance, Eingreifen bei Abweichung. Kontinuierliche Aufsichtsarchitektur ersetzt periodische Reviews; Rekalibrierung wird signalgetrieben.

Gefährlichstes Risiko: L4-Autonomie auszurollen, bevor die menschliche Fähigkeit existiert, sie zu orchestrieren.

„Organisationen, die oberhalb dieser Schwelle operieren wollen — Stufe-3- oder Stufe-4-Entscheidungsautorität auf Stufe-1-Organisationsinfrastruktur —, beschleunigen ihre Transformation nicht. Sie akkumulieren Governance-Schulden, die irgendwann eine Abrechnung erzwingen — typischerweise zum denkbar schlechtesten Zeitpunkt.“

— ÜBER DAS VERHÄLTNISS VON REIFE UND AUTONOMIE

ABSCHNITT 11

Warum Integration entscheidend ist.

Jede der sechs Domänen produziert isoliert betrachtet ein unvollständiges und letztlich nicht tragfähiges Ergebnis. Integration schafft das organisatorische System, das zu nachhaltiger KI-Wertschöpfung fähig ist.

GOVERNANCE & STRUKTUR

Eine Organisation, die eine ausgefeilte Governance-Architektur baut, ohne die Center-Led-Struktur zu bauen, die die organisatorische Kapazität für gelebte Governance schafft, wird ein Governance-Framework haben, das auf dem Papier existiert — und ein produktives KI-Portfolio, das ohne echte Aufsicht operiert. **Governance braucht das organisatorische Substrat klarer Ownership und die operative Bandbreite, diese Ownership auszuüben.** Struktur liefert beides.

PROZESS & FÄHIGKEITEN

Eine Organisation, die stark in AI-Integrated Process Redesign investiert, ohne komplementäre rollenspezifische Fähigkeitsentwicklung, wird feststellen, dass die neuen Prozesse nicht wie entworfen funktionieren. Übergabekriterien werden falsch angewendet. Eskalationsprotokolle werden ignoriert. Ausnahmebehandlung fällt in informelle Arrangements zurück. **Prozessdesign sagt Menschen, wie Arbeit fließen soll. Fähigkeitsentwicklung entscheidet, ob sie sie tatsächlich so fließen lassen können.**

STRATEGIE & DECISION ARCHITECTURE

Eine Organisation, die eine Capability-Based Investment Thesis definiert — etwa Customer Intelligence als strategische KI-Fähigkeit —, ohne zugleich die Decision Rights Registry zu bauen, die formal autorisiert, welche kundengerichteten Entscheidungen KI auf welcher Autonomiestufe ausführen darf, wird feststellen: Die Fähigkeitsinvestition produziert KI-Systeme, auf deren Ausgaben niemand formal zu handeln befugt ist.

EIN DESIGNSYSTEM, KEINE CHECKLISTE

Diese Interdependenzen lassen sich nicht als sequenzielle Phasen eines Umsetzungsplans managen. Sie müssen als **Designsystem** gemanagt werden: ein Set von Entscheidungen, die in Kenntnis voneinander getroffen, iterativ justiert und durch den gemeinsamen Referenzpunkt des Human-AI Interface zusammengehalten werden. **Das bedeutet es, eine Organisation zu gestalten, statt ein Projekt zu managen.**

FAZIT

KI transformiert keine Organisationen. Organisationen transformieren sich selbst, um mit KI zu arbeiten.

Das Argument dieses Whitepapers ist letztlich einfach, auch wenn seine Ausarbeitung komplex ist. Die Transformation erfordert Design — nicht Projektmanagement, nicht Change Management, nicht KI-Strategie in Isolation, sondern echte Organisationsdesign-Arbeit, die Struktur, Governance, Entscheidungsfindung, Prozesse und Fähigkeiten als integriertes System adressiert.

Das HandsOn AI Operating Model liefert das Framework für diese Designarbeit. Es schöpft aus den rigorosesten Traditionen der Organisationsdesign-Methodik und erweitert sie bewusst in den KI-nativen Kontext. Es stellt das Human-AI Interface — das explizite Design des Verhältnisses von menschlichem Urteil und KI-Handlungsfähigkeit — ins Zentrum jeder Domäne. So sind die sechs Elemente nicht nur aufeinander ausgerichtet, sondern auf die fundamentale organisatorische Frage des KI-Zeitalters.

Das Zeitfenster, diese Arbeit bewusst statt reaktiv zu leisten, ist nicht unbegrenzt. Organisationen, die die Operating-Model-Infrastruktur für Stufe-2- und Stufe-3-Deployment jetzt aufbauen, werden diesen Vorteil über die Zeit verzinsen. Organisationen, die weiter pilotieren und augmentieren, ohne neu zu gestalten, werden erleben, dass sich die Lücke zwischen ihren KI-Investitionen und ihrem KI-Wert vergrößert statt schließt.

Der Engpass war nie die Technologie. Der Engpass ist und bleibt die Organisation. Die Arbeit, diese Lücke zu schließen, ist die Arbeit der KI-Transformation — und sie beginnt mit Design.

„Bereit, das Operating Model für KI zu bauen, das wirklich funktioniert?“

— MIT HANDSON ARBEITEN

ZUSAMMENARBEIT

Drei Wege, mit HandsOn zu arbeiten.

HandsOn gestaltet und implementiert AI Operating Models auf Basis von Organisationsdesign-Methodik. Wir arbeiten mit Führungsteams in dem Moment, in dem die Pilotphase endet und die harte organisatorische Arbeit beginnt.

01

AI Operating Model Diagnostic

Ein strukturiertes 4–6-wöchiges Assessment aller sechs Design-Domänen. Geliefert als Reifelandkarte, Gap-Analyse und priorisierte Design-Agenda für Ihr Führungsteam.

DAUER · 4–6 WOCHEN

02

AI Operating Model Design Sprint

Ein intensives 8–12-wöchiges Engagement, um Ihr Ziel-Operating-Model über die sechs Domänen zu gestalten — inklusive Governance-Architektur, Entscheidungsrechte-Framework und Capability-Roadmap.

DAUER · 8–12 WOCHEN

03

Full AI Transformation Partnership

Ende-zu-Ende-Begleitung von der Diagnose bis zur Implementierung — HandsOn-Expertise direkt eingebettet in Ihre Transformationsteams für nachhaltigen organisatorischen Wandel.

DAUER · FORTLAUFEND

ÜBER DEN AUTOR

Maximilian Stein — Gründer, HandsOn. Zuvor VP Marketing & Communication bei FERCHAU. Maximilian hat große Transformationen erfolgreich geführt und verbindet Strategie, Operating Models und Umsetzung, um Ergebnisse zu liefern.

ÜBER HANDSON

HandsOn ist auf die organisatorische Transformation spezialisiert, die Unternehmen brauchen, damit KI wirklich funktioniert: der Aufbau von AI Operating Models, KI-Governance-Architekturen und KI-Fähigkeitsprogrammen — fundiert in Organisationsdesign-Methodik und kalibriert auf den spezifischen operativen Kontext jedes Klienten.

KONTAKT

wearehandson.de

Für Diagnostik-Engagements, Design Sprints und laufende Transformationspartnerschaften erreichen Sie uns über die obige Adresse.

© 2026 HandsOn — AI Strategy & Governance Consulting. Alle Rechte vorbehalten. The HandsOn AI Operating Model™ ist eine Marke von HandsOn. Dieses Dokument ist für die Weitergabe an Klienten bestimmt und darf vollständig reproduziert werden, sofern die Quellenangabe erhalten bleibt. Im Text zitierte Quellen: Deloitte State of Generative AI 2026 · Deloitte AI Institute 2026 · McKinsey State of AI 2025.