

KI und das Lagebild – Vom Bericht zum Lagedienst

Wie Künstliche Intelligenz Erstellung und Bereitstellung von Lageinformationen verändert

Version 2.1 · Stand 18.08.2026 · Zielgruppe: Entscheider einer Enterprise-Unternehmung – KI-, Militär- oder Agile-Vorwissen wird nicht vorausgesetzt

Zu diesem Dokument: Kuratiert, inhaltlich durchdacht und verantwortet von **Robert Seebauer**; recherchiert und erstellt mit Unterstützung von **Claude** (Anthropic – Modell **Claude Fable 5**, via Claude Code). Dies ist **Version 2.1**: die Vertiefung der v2.0 nach Lektüre und Bewertung von *Prediction Machines*, *Noise* und *Co-Intelligence* durch den Kurator (**Teil A2 als „Urteilkraft-Block“**, Kapitel **C.6, Entscheidungshygiene** als siebtes KI-Prinzip, erweitertes Kompetenzfeld 7, Noise-Audit im Pilot, Roadmap-Konsequenzen für Phase D) – ergänzt um den **Rechtsrahmen der KI-Verordnung** auf dem Stand August 2026 (Rechts-Prüfbox in D.2: Betreiberpflichten Art. 4 und Art. 50, Hochrisiko-Fälle Personenbezug, rechtliche Leitplanken für Phase D) und um den Transparenzhinweis nach dem Reihen-Standard. Schwesterdokumente: „*Team of Teams & Stabsausbildung*“ (Menschen & Rituale) und „*Enterprise-Architektur & Solution-Lagebild*“ (Inhalt & Struktur) – zusammen mit dieser Schrift (Werkzeug & Bereitstellung) der Dreiklang der Reihe. Ältere Stände: [Release-Archiv](#) [denkschriften/<Datum>](#) .

Transparenzhinweis: Diese Schrift wurde mit Unterstützung eines KI-Systems (Claude, Anthropic – Modell Claude Fable 5, via Claude Code) erstellt: Entwurf, Recherche-Aufbereitung, Satz und Grafik-Code. Alle Inhalte wurden von Robert Seebauer inhaltlich geprüft, redaktionell bearbeitet und freigegeben; er trägt die redaktionelle Verantwortung.

Executive Summary

Ausgangsfrage: Die Reihe hat gezeigt, *wie* ein Lagebild geführt wird (Rituale, Stabsausbildung) und *was* hineingehört (Capabilities, Schulden, Konfidenz). Offen ist die Werkzeugfrage: **Inwieweit kann Künstliche Intelligenz beim Erstellen eines Lagebilds und – vor allem – bei der Bereitstellung von Lageinformationen helfen?**

WICHTIGSTER EINZELSATZ DES DOKUMENTS

KI macht Lageinformation billig, Urteilkraft bleibt knapp – deshalb gehört KI in den Stab, nicht auf den Feldherrnhügel, und ihr größter Hebel ist die **Bereitstellung**: die richtige Information, zur richtigen Zeit, in der Sprache des Empfängers.

Befund in fünf Sätzen:

- 1 Die Ökonomie dreht sich.** KI senkt die Kosten von Informationsaufbereitung und Vorhersage drastisch; im Gegenzug steigt der Wert des komplementären Gutes – menschlicher Urteilkraft (Agrawal/Gans/Goldfarb). Der Engpass des Lagebilds wandert von der *Erstellung* zur *Deutung und Zustellung*.
- 2 Das Militär liefert das reife Modell.** Der Nachrichtenzyklus behandelt Bereitstellung („Dissemination“) seit Jahrzehnten als *eigene Phase mit eigener Doktrin* – und die KI-Praxis von NATO und Streitkräften zeigt eine klare Linie: KI verdichtet Lage und Analyse, Entscheidung und Verantwortung bleiben beim Menschen.

- 3 **Der größte zivile Hebel ist der Lagedienst.** KI macht erstmals einen *stehenden, befragbaren* Lagedienst fürs Portfolio bezahlbar – den Wandel vom periodischen Bericht zum Dienst, den man jederzeit befragt und der sich an Entscheidungspunkten selbst meldet. Das Security Operations Center beweist seit Jahren, dass Unternehmen so etwas betreiben können.
- 4 **Ohne Herkunft keine Wahrheit – und ohne Hygiene kein Urteil** (*vertieft in v2.0*). Ein KI-generiertes Lagebild ohne Quellen- und Konfidenzausweis ist die perfekte „Information Masquerade“ (Milchman). Und Urteile leiden nicht nur an Verzerrung (Bias), sondern an **Streuung (Noise)** – KI kann beides dämpfen, verstärkt aber Noise, wenn ihr blind vertraut wird (Kahneman/Sibony/Sunstein). Die Empirie bestätigt: ~90 % KI-Adoption, ~30 % Misstrauen, und KI *verstärkt*, was da ist (DORA 2025). **Sieben Prinzipien** ziehen daraus die Leitplanken – neu darunter die Entscheidungshygiene.
- 5 **Konsequenz für uns:** Jede Entscheidung lässt sich in sechs Bestandteile zerlegen (Agrawal/Gans/Goldfarb) – Kapitel C.6 ordnet je Bestandteil zu, wo unsere Software heute hilft, wo KI helfen kann und wo der Mensch bleibt. Dazu Kompetenzfeld 7 und Curriculum-Modul 7, ein 90-Tage-Pilot mit Delta-Briefing und **Noise-Audit**, und die Software-Roadmap **Phase D** hinter dem deterministischen Kern.

Was hier „KI“ heißt (Taxonomie, gilt für das ganze Dokument)

Stufe	Was	Bei uns heute
1 · Regelbasierte Automatisierung	Deterministische Pipelines, Schwellwerte, Ampeln	Gesamte Report-Pipeline
2 · Statistik & klassisches ML	Monte-Carlo-Forecasts, Anomalieerkennung, Fehlervorhersage	Simulate-Modul; <i>Fault Detection</i>
3 · Generative KI / LLMs (Schwerpunkt)	Zusammenfassen, Übersetzen (Sprachen und Lesarten), Q&A, Entwürfe	Denkschriften-Erstellung, Manuals in 5 Sprachen
4 · Agentische KI	Selbständige Erhebung, mehrstufige Aufträge, Werkzeugnutzung	Nur Ausblick (D.2, Guardrails)

Wer über „KI im Lagebild“ entscheidet, sollte immer sagen können, über welche Stufe er gerade spricht – Chancen, Risiken und Leitplanken unterscheiden sich je Stufe erheblich.

Teil A – Die Quellen: KI im eigenen Korpus

Die vier KI-Titel des 22-bändigen Apress-Korpus (Volltexte im Werkstatt-Archiv des Kurators), gelesen mit der Lagebild-Brille – plus Milchman als Scharnier zur EA-Denkschrift.

A.1 *Architecting Enterprise AI Applications* – die strategische Grenzziehung

Das strategisch-konzeptionelle Werk des Korpus zur Frage, **wo KI entscheiden sollte und wo nicht**: Mensch-vs.-KI-Aufgabenteilung, Meta-Systeme, Overfitting- und Proxy-Fallen. Kernhaltung: KI bewusst *begrenzt* einsetzen, wo Kontext und Verantwortung zählen.

Für uns: der Korpus-Kronzeuge für das Rollenbild dieser Schrift – KI als Stabsmitglied mit definiertem Auftrag, nicht als Entscheider.

A.2 *Mastering Machine Learning Architecture and Solutions* – der Lifecycle

ML-Systeme als Architektur-Bausteine mit eigenem **Lebenszyklus**: Datenpipelines, Modellwahl, Betrieb, Drift – MLOps als „DevOps für Modelle“.

Für uns: Wer KI ins Lagebild einbaut, übernimmt Betriebsverantwortung (Risikokatalog D.2: Modell-Drift); KI-Bausteine sind keine Einmal-Installation.

A.3 *AI and Microservices* – die angewandte Sicht

KI-gestützte Entwicklung und API-/Testarbeit in Microservice-Landschaften.

Für uns: Beleg, dass Stufe-3-KI im Engineering-Alltag bereits Werkzeugstatus hat – die Frage ist nicht *ob*, sondern *mit welchen Leitplanken* (deckungsgleich mit dem DORA-Befund in A2.6).

A.4 *Fault Detection in Microservice Architectures* – Auswerten in Reinform

Datengetriebene **Fehlervorhersage** aus Struktur-Metriken (Coupling/Cohesion → Fehlerwahrscheinlichkeit, bewertet mit Precision/Recall).

Für uns: das konkrete Muster, wie aus vorhandenen Metriken *Frühindikatoren* werden – Stufe-2-KI für die Auswerten-Phase des Zyklus (C.1), lange bevor ein LLM ins Spiel kommt.

A.5 Das Scharnier: Milchman, *Unit Oriented Enterprise Architecture*

Das einzige Korpus-Werk, das wörtlich von „**situational awareness**“ spricht – Untertitel: *Constructing Large Sociotechnical Systems in the Age of AI*. Seine Diagnose der „**Information Masquerade**“ (Information fließt schlechter, als alle behaupten) wird im KI-Zeitalter zur doppelten Warnung: KI kann die Maskerade *bekämpfen* (Herkunft, Konfidenz, Verfügbarkeit) oder *perfektionieren* (flüssiger Text ohne Substanz). Welche Seite gewinnt, entscheiden die Prinzipien in C.4.

Teil A2 – Der Urteilskraft-Block: Ökonomie, Streuung, gezackte Grenze – und die Automatisierungsforschung

Sieben Quellen außerhalb des Korpus, die je eine Flanke tragen. Die ersten drei – in v1.x nur nach allgemeinem Kenntnisstand skizziert – hat der Kurator inzwischen vollständig gelesen und bewertet; sie bilden gemeinsam mit Duke (A2.7) den **Urteilskraft-Block** dieser Denkschrift: Warum Urteilskraft das knappe Gut ist (Prediction Machines), woran sie

krankt (Noise), wie man KI einbindet, ohne sie zu verspielen (Co-Intelligence), und wie man sie trainiert (Duke). Jede Würdigung endet mit der **Position des Kurators** – der Stelle, an der aus Lektüre eine Haltung wird.

A2.1 Agrawal/Gans/Goldfarb – *Prediction Machines* (2018, akt. Aufl. 2022): die Anatomie einer Entscheidung

Die These der Rotman-Ökonomen: KI ist ökonomisch betrachtet ein **Preissturz für Vorhersage** – verstanden im weiten Sinn als das Ausfüllen fehlender Information. Wird ein Gut billig, steigt der Wert seiner **Komplemente**: Daten und vor allem **Urteilkraft** (die Bewertung, was eine Vorhersage *bedeutet* und welches Ergebnis wie viel wert ist). Das trägt den Wichtigsten Einzelsatz dieser Schrift – doch das eigentliche Arbeitswerkzeug des Buches ist die **Zerlegung jeder Aufgabe in sechs Bestandteile** (der „AI Canvas“):

Bestandteil	Frage	Wer/was leistet es
Vorhersage	Was werden wir wissen, wenn wir es wissen? Welche fehlende Information wird ergänzt?	die Domäne der Maschine – hier fällt der Preis
Urteil	Wie bewerten wir die möglichen Ergebnisse? Was ist uns Treffer und Fehler jeweils wert?	menschlich; Komplement, das im Wert steigt
Aktion	Was tun wir – welche Handlung folgt aus Vorhersage plus Urteil?	menschlich oder automatisiert – je nach Regler (C.2)
Ergebnis	Woran messen wir, ob die Aktion erfolgreich war?	die Rückmeldung an alle anderen Bestandteile
Input-Daten	Welche Daten braucht die Vorhersage im laufenden Betrieb?	Voraussetzung – hier entscheidet sich Qualität
Training/Feedback	Aus welchen Daten lernt das System, und wie fließen Ergebnisse zurück?	der Lernkreis – ohne ihn veraltet jede Vorhersage

Die Kraft der Zerlegung liegt darin, dass sie **die Frage „Wo hilft KI?“ beantwortbar macht**: nicht pauschal, sondern je Bestandteil. Und sie enthält eine ökonomische Warnung: Wer Vorhersage und Urteil verwechselt – also der Maschine überlässt, was Ergebnisse *wert* sind –, hat nicht automatisiert, sondern abgedankt.

POSITION DES KURATORS

Volle Zustimmung. Die Zerlegung ist schlüssig und dienlich – sie ist der Grund, warum diese Schrift in v2.0 ein eigenes Kapitel bekommt (C.6), das je Bestandteil prüft, wo genau unsere Software heute hilft, wo KI helfen kann und wo der Mensch bleibt.

A2.2 Kahneman/Sibony/Sunstein – *Noise* (2021): Urteile streuen – und KI kann beides

Das Buch macht eine Unterscheidung, die Entscheider meist nicht treffen: **Bias** ist die *systematische* Abweichung von Urteilen (alle liegen in dieselbe Richtung daneben) – **Noise** ist ihre *unerwünschte Streuung* (dieselbe Lage, dieselben Fakten, unterschiedliche Urteile – je nach Person, Tagesform, Reihenfolge, Wetter). Die Autoren zeigen, dass Noise in Organisationen fast überall groß, fast nie gemessen und fast immer unterschätzt ist – von Versicherungsprämien über Gerichtsurteile bis zu Personalentscheidungen. Beide Fehlerquellen sind unabhängig voneinander; **wer nur Bias bekämpft, lässt die Hälfte des Fehlers liegen**. Ihre Antwort ist die **Entscheidungshygiene** – Regeln, die Noise reduzieren, ohne dass man ihn im Einzelfall sehen muss:

- **Unabhängige Urteile zuerst, Diskussion danach** – wer die Meinung des Ranghöchsten oder der Gruppe kennt, urteilt nicht mehr unabhängig.
- **Aggregation** – der Durchschnitt mehrerer unabhängiger Urteile ist zuverlässiger als das einzelne, gerade bei Expertise.
- **Strukturierte Bewertung** – komplexe Urteile in Teilurteile zerlegen, jedes einzeln bewerten, das Gesamturteil bis zum Schluss zurückhalten.
- **Relative statt absolute Skalen** – Ranking und Vergleich sind konsistenter als absolute Noten.
- **Noise-Audits** – dieselbe Lage mehreren Beurteilern getrennt vorlegen und die Streuung messen; erst dann weiß eine Organisation, wie laut sie ist.
- **Algorithmen und Regeln als Konsistenz-Anker** – ein einfaches, transparentes Modell schlägt oft die Streuung menschlicher Urteile, weil es nicht müde wird und keine Tagesform hat.

Der letzte Punkt macht *Noise* für diese Schrift ambivalent: KI kann Noise **dämpfen** (immer dieselbe Regel, dieselbe Aufmerksamkeit) – aber ein *großes Sprachmodell* ist kein einfaches Modell, sondern selbst **nicht deterministisch** und mit **eigenen Fehlerquellen** (Halluzination, falsche Quellen, Sensitivität für die Formulierung der Frage). Wer ihm blind vertraut, importiert dessen Streuung ins eigene Urteil – und nimmt sie wegen der Objektivitäts-Anmutung nicht einmal mehr wahr.

POSITION DES KURATORS

Volle Zustimmung – Bias und Noise sind verschieden, und man muss immer **beide** betrachten. Für KI im Lagebild überwiegt die Gefahr: Ich sehe KI eher als **Noise-Verstärker**, weil ihr blind vertraut wird, ohne kritisch zu hinterfragen – während sie halluziniert oder auf falschen Quellen aufbaut. Die Entscheidungshygiene-Regeln sind wichtig und werden übernommen: als siebtes Prinzip (C.4), in Kompetenzfeld 7 (C.5), als Noise-Audit im Piloten (D.3).

A2.3 Mollick – Co-Intelligence (2024): die gezackte Grenze und der Mensch in der Schleife

Der Wharton-Professor destilliert den Arbeitsalltag mit Sprachmodellen zu vier Regeln: **KI immer mit an den Tisch holen** (probieren, was sie in der eigenen Arbeit kann), **selbst der Mensch in der Schleife bleiben, die KI wie eine Person behandeln – aber ihr sagen, welche** (Rolle, Kontext, Auftrag), und **annehmen, dass dies die schlechteste KI ist, die man je benutzen wird** (die Modelle entwickeln sich schneller, als Organisationen lernen). Sein wichtigster Befund ist die „**jagged frontier**“: Die Fähigkeitsgrenze der Modelle verläuft gezackt – überraschend stark bei Aufgaben, die schwer aussehen, überraschend schwach bei solchen, die leicht aussehen –, und diese Grenze ist **von außen unsichtbar**. Man findet sie nur durch Benutzung, und man findet sie an jeder Stelle neu.

POSITION DES KURATORS

Das Konzept ist schlüssig und deckt sich mit meinen eigenen Beobachtungen – gerade *weil* KI nicht deterministisch ist. Daraus folgt kein Zögern, sondern eine doppelte Konsequenz: **KI stark einbinden, um ihre Grenzen überhaupt zu erkennen** – aber das **immer mit dem Menschen in der Schleife**, denn wir sehen Grenzen und Glanzpunkte heute noch gar nicht wirklich, und die Modelle entwickeln sich zu schnell, um zu gamblen. Ausprobieren ist Pflicht; Vertrauen ohne Prüfung ist Leichtsinn.

A2.4 Bainbridge – „Ironies of Automation“ (1983): die Gegenrede

Der Klassiker der Automatisierungsforschung (doi.org), in vier Seiten: Automatisierung übernimmt die *leichten* Teile einer Aufgabe und lässt dem Menschen die *schwersten* – Überwachung und Ausnahmebehandlung –, während sie ihm zugleich die tägliche Übung entzieht, die er genau dafür bräuchte. Je besser die Automatisierung, desto **wichtiger und zugleich schlechter vorbereitet** der Mensch.

Für uns: die vielleicht wichtigste Einzelwarnung dieser Schrift – Prinzip 5 „Der Mensch übt weiter“ (C.4) und die KI-freien Katas im Curriculum (D.4) sind die direkte Antwort.

A2.5 Parasuraman/Sheridan/Wickens (2000): der Regler

Das Standardmodell der Mensch-Automation-Forschung unterscheidet **vier Funktionsklassen** – Informationsgewinnung, Informationsanalyse, Entscheidungsauswahl, Ausführung – und je Klasse **Automatisierungsgrade** von „Mensch macht alles“ bis „Maschine entscheidet autonom“. Die Pointe: Der Automatisierungsgrad ist **je Funktionsklasse getrennt zu wählen**.

Für uns: die wissenschaftliche Grundlage des Automatisierungsreglers in C.2 – „Stabsmitglied, nicht Kommandeur“ wird zur Stellregel statt zum Verbot.

A2.6 DORA 2025 – State of AI-assisted Software Development: die Empirie

Die jährliche DORA-Studie (Google Cloud, ~5.000 Befragte weltweit) liefert drei Befunde für Entscheider: (1) **~90 % der Software-Fachleute nutzen KI** bei der Arbeit – die Adoptionsfrage ist durch. (2) Der „**Amplifier-Effekt**“: KI verstärkt vorhandene Stärken und Schwächen; leistungsfähige Organisationen werden besser, schwache sehen ihre Probleme deutlicher – KI *repariert* nichts. (3) Trotz breiter Nutzung berichten **~30 % wenig oder kein Vertrauen** in KI-generierten Output; den größten Ertrag bringen nicht die Tools, sondern Plattform-Qualität und klare Workflows.

Für uns: (2) ist die empirische Bestätigung des ToT-Einzelsatzes („Ritual vor Tool“) auf KI übertragen; (3) macht Herkunft und Konfidenz zur Akzeptanz-, nicht zur Stilfrage.

A2.7 Duke – *Thinking in Bets* (2018): die Urteils-Hygiene (seit v1.1)

Die frühere Poker-Profi behandelt das Bewerten unter Unsicherheit selbst – drei Beiträge tragen direkt: **Resulting** (Entscheidungs- von Ergebnisqualität trennen) ist die Lesekompetenz, ohne die probabilistische Lageinformation politisch nicht überlebt – ein eingetretenes 15-%-Szenario widerlegt keine P85-Prognose. **Truthseeking-Gruppen** (Wahrheitssuche-Charta: Information teilen, organisierte Skepsis, Rechenschaft) liefern das *soziale* Gegenmittel gegen Automation Bias – die Gruppen-Gegenprobe, bevor eine KI-Vorlage übernommen wird. **Ulysses-Verträge** (Vorab-Festlegungen, die das spätere Ich binden) geben dem Entscheidungspunkt-Wecker (M4) sein Fundament.

Für uns: Duke trainiert genau die Urteilskraft, die nach *Prediction Machines* das knappe Gut ist – und ihre Truthseeking-Gruppen sind die soziale Form der Entscheidungshygiene aus *Noise*: unabhängig urteilen, dann aggregieren, dann diskutieren. (Ausführliche Würdigung: ToT-Denkschrift, Teil A2.5.)

DER URTEILSKRAFT-BLOCK AUF EINEN SATZ

Urteilskraft ist das knappe Gut (*Prediction Machines*), sie streut mehr, als wir glauben (*Noise*), KI findet ihre eigenen Grenzen nur mit dem Menschen in der Schleife (*Co-Intelligence*) – und trainierbar ist sie durch Wetten, Truthseeking und die Trennung von Entscheidung und Ergebnis (Duke).

Teil B – Referenzpraxis: das Nachrichtenwesen und seine KI

Die Methode der Reihe: erst die reife Fremddisziplin verstehen, dann übertragen. Für Lageinformation ist das reifste System das militärische Nachrichtenwesen – es beantwortet unsere Leitfrage seit Jahrzehnten institutionell.

B.1 Der Nachrichtenzyklus im Klartext

Die US-Doktrin (Joint Publication 2-0, *Joint Intelligence*) organisiert Nachrichtenarbeit als Kreislauf aus

Planung/Steuerung → Beschaffung → Aufbereitung → Auswertung → Verteilung/Integration, mit ständiger Bewertung und Rückkopplung. Drei Eigenschaften machen ihn für uns wertvoll:

1 Er beginnt beim Entscheider, nicht bei den Daten

Am Anfang steht der priorisierte Informationsbedarf des Führers (**CCIR/PIR** – Commander's Critical / Priority Intelligence Requirements): *Welche Information würde meine Entscheidung ändern?* Das ist exakt Hubbards EVI-Frage in doktrinäer Form – erhoben wird, was Entscheidungen ändern kann, nicht was messbar ist.

2 Verteilung ist eine eigene Phase mit eigener Doktrin

Die beste Auswertung ist wertlos, wenn sie den Entscheider nicht rechtzeitig, verständlich und im richtigen Format erreicht – darum regelt die Doktrin ausdrücklich *Push* (vorausgeplante Zustellung an definierte Empfänger) und *Pull* (Zugriff auf Bestände nach Bedarf), Empfängergerechtigkeit und Aktualität. **Genau diese Phase adressiert unsere Leitfrage mit „vor allem der Bereitstellung“.**

3 Er ist ein Kreislauf

Jede Entscheidung erzeugt neuen Informationsbedarf; Bewertung läuft ständig mit. Ein Lagebild ist nie „fertig“.

B.2 Was die Doktrin übers Bereitstellen lehrt

Aus der Verteilungs-Doktrin und der Stabspraxis (vgl. ToT-Denkschrift, Teil B) lassen sich fünf Bereitstellungs-Lehren destillieren, die unabhängig von KI gelten – KI macht sie nur bezahlbar:

- **BLUF – Bottom Line Up Front:** Kernaussage zuerst, Herleitung danach; der Lagevortrag ist eine geübte Form, kein Prosastück.
- **Empfängergerecht:** Derselbe Sachverhalt wird je Ebene anders aufbereitet (Kommandeur ↔ Fachstab) – *ohne* dass sich die Fakten ändern.
- **Push nach Plan, Pull nach Bedarf:** Kritisches wird aktiv zugestellt (an Entscheidungspunkten), alles andere ist auf Abruf verfügbar.
- **Aktualität vor Perfektion:** Ein rechtzeitiges 80-%-Bild schlägt ein perfektes von gestern (McChrystals O&I lebte von diesem Prinzip).
- **Herkunft und Verlässlichkeit gehören zur Meldung:** Das Nachrichtenwesen bewertet Quellen seit jeher explizit – die zivile Entsprechung sind unsere Konfidenz-Flags (EA-Prinzip 4).

B.3 KI im Bündnis: was NATO und Forschung berichten – und was beim Menschen bleibt

- **NATO ACT, Decision-Making in the Age of Big Data and AI** (2021): Multimodale Sensordaten plus Analytik sollen ein **konsolidiertes, mehrdimensionales Lagebild nahezu in Echtzeit** ermöglichen; KI entlastet Analysten von Sichtungs- und Fusionsarbeit und beschleunigt Entscheidungszyklen. Ebenso klar benennt der Bericht die Voraussetzungen: Datenqualität, Interoperabilität, Vertrauen der Nutzer – und menschliche Verantwortung.
- **RAND (RRA4316-1)** analysiert KI entlang vier militärischer Grundwettbewerbe – Menge vs. Qualität, **Verbergen vs. Finden, zentrale vs. dezentrale Führung**, Cyber-Angriff vs. -Verteidigung. Für uns sind zwei zentral: KI ist primär eine *Finde-Technologie* (sie hebt Signale aus Rauschen – das ist Lagebild-Arbeit), und sie kann **beide** Führungsmodelle

stärken. Das ist eine Warnung: Wer mit KI alles zentral sehen *kann*, wird versucht sein, wieder alles zentral zu *entscheiden* – McChrystals Kopplung (Lagebild **plus** dezentrale Befugnis) muss gegen diesen Rezentralisierungs-Reflex aktiv verteidigt werden.

- **Atlantic Council** (Bericht zur KI-Integration der NATO): KI als Enabler konsolidierter Multi-Domain-Lagebilder in C4ISR – dieselbe Stoßrichtung aus Bündnis-Perspektive.
- **Die Grenze zieht die NATO selbst:** Ihre *Principles of Responsible Use* (2021) verlangen u. a. Verantwortlichkeit (Menschen bleiben rechenschaftspflichtig), Nachvollziehbarkeit und Steuerbarkeit von KI-Einsätzen. Übersetzt in unser Bild: **KI dient im Stab, sie führt nicht** – die Menschen bleiben es, die sie steuern und die Entscheidung verantworten. Denselben Grundsatz schreibt der europäische Gesetzgeber fest: Die **KI-Verordnung** verlangt vom *Betreiber* eines KI-Systems Kompetenz (Art. 4) und Transparenz (Art. 50) und lässt Verantwortung und Urteil beim Menschen – die Konsequenzen für den Lagedienst stehen in der Rechts-Prüfbox in D.2.

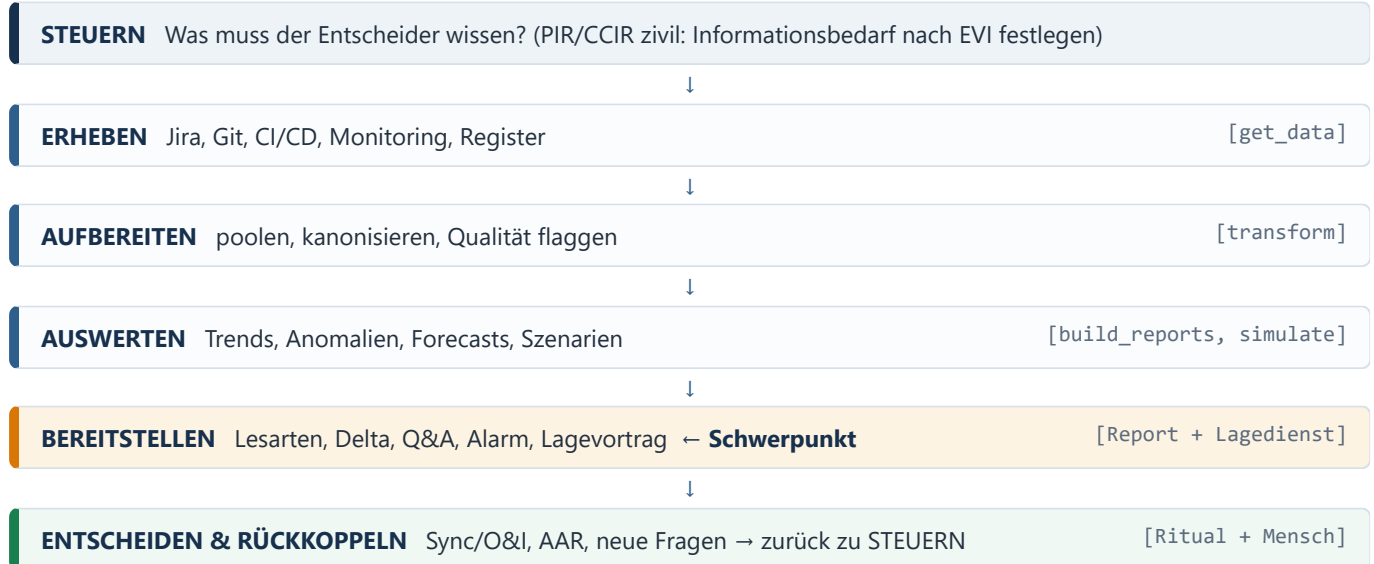
B.4 Der zivile Existenzbeweis: das Security Operations Center

Unternehmen betreiben längst einen stehenden Lagedienst – dort, wo es brennt: Das **SOC** hält rund um die Uhr ein Cyber-Lagebild. Sensorik liefert Ereignisse, Korrelationssysteme (SIEM) verdichten sie, **KI-gestützte Triage** priorisiert, Alarme kommen mit Kontext, und ein menschlicher Analyst entscheidet über die Reaktion. Zwei Lektionen aus zwanzig Jahren SOC-Praxis übertragen sich direkt: **Alarm-Müdigkeit** ist der Tod jedes Lagedienstes (darum Priorisierung nach Entscheidungsrelevanz – EVI), und **Automatisierung endet vor der Reaktion** – auch im SOC löst die Maschine selten selbst aus. Für das Portfolio, wo über Millionen entschieden wird, existiert nichts Vergleichbares – nicht mangels Bedarf, sondern weil ein menschlich besetzter Lagedienst unbezahlbar wäre. **Diese Rechnung ändert KI.**

Teil C – Übertragung: der KI-gestützte Lagedienst fürs Unternehmen

C.1 Der Zyklus fürs Enterprise-Lagebild – und die vier Rollen der KI

Der Nachrichtenzyklus mappt direkt auf die Lagebild-Arbeit der Solution-/Portfolio-Ebene (und auf unsere Software-Pipeline):



Quer zum Zyklus lässt sich der KI-Beitrag in **vier Rollen** fassen – alle vier sind Stabsrollen, keine Führungsrollen:

Rolle	Zyklus-Schritte	Beispiele	Reifegrad
1 · Sensor & Sammler	Erheben, Aufbereiten	Extraktion, Klassifikation („welche Capability betrifft dieses Epic?“), Dubletten, Qualitäts-Flags begründen	heute produktiv
2 · Stabsgehilfe (Redakteur)	Aufbereiten, Bereitstellen	Zusammenfassungen, Lesarten-Übersetzung, Delta-Briefings, Lagevortrag-Entwurf, Mehrsprachigkeit	heute produktiv, mit Prüfpflicht
3 · Analyst & Sparringspartner	Auswerten	Anomalie-/Trenderkennung, Forecasts (MC/ML), Szenario-Entwürfe fürs Wargaming, Premortem-Fragen, Annahmen-Angriff (Red Team)	teils produktiv, teils erprobenswert
4 · Lagedienst	Bereitstellen	Befragbares Lagebild, Entscheidungspunkt-Wecker, adressatengerechte Zustellung	erprobenswert – das Neuland

C.2 Der Automatisierungsregler: wie viel KI je Schritt?

Nach Parasuraman/Sheridan/Wickens (A2.5) wird der Automatisierungsgrad **je Funktionsklasse** eingestellt – für das Lagebild empfehlen wir vier Reglerstände:

- 1 Informationsgewinnung und -aufbereitung: Regler hoch**
Hier entsteht die Kostensenkung; menschliche Fleißarbeit in Datensammlung ist kein Qualitätsmerkmal, sondern Verschwendung.
- 2 Analyse: Regler mittel bis hoch**
Mit Konfidenzausweis an jedem Ergebnis und Stichproben-Verifikation als Stabsroutine.

3 Entscheidungsauswahl: Regler niedrig

KI darf Optionen entwerfen, vergleichen und empfehlen – gewählt wird von Menschen (das ist zugleich die MDMP-Logik: der Stab entwickelt Optionen, der Kommandeur entscheidet).

4 Ausführung: nur mit menschlicher Freigabe

Im Lagebild-Kontext heißt das praktisch: keine automatischen Eskalationen oder Umpriorisierungen ohne vereinbarte Regel und benannten Verantwortlichen.

Damit wird „KI ist Stabsmitglied, nicht Kommandeur“ von der Metapher zur Stellregel – differenziert genug für echte Architekturentscheidungen, einfach genug für die Vorstandsfolie.

C.3 Der Lagedienst: sieben Bereitstellungs-Muster (das Herzstück)

Der Schwerpunkt der Leitfrage liegt auf der Bereitstellung – hier sind die sieben Muster, mit denen KI aus dem Bericht einen Lagedienst macht:

#	Muster	Was es leistet	Hauptrisiko	Leitplanke
M1	Adressaten-Lesarten	Executive-BLUF und Engineering-Detail aus denselben Zahlen; erweitert EA-Prinzip 3 zu „ein Bild, n Lesarten“	Personalisierung zersplittert die eine Wahrheit	Eine Datenbasis, gemeinsamer Kennzahlen-Kern, Provenienz je Aussage
M2	Delta-Briefing	„Was hat sich seit dem letzten Sync geändert – und was beschleunigt sich?“ (Hohpes „first derivative“ als Produkt)	Rauschen wird als Änderung verkauft	Deterministischer Diff liefert die Fakten, KI nur die Formulierung
M3	Befragbares Lagebild	Q&A in Alltagssprache über die Report-Daten; Drill-down ohne Werkzeughürde – mehr Menschen können das Lagebild <i>lesen</i>	Halluzination: eloquente Antworten ohne Basis	Zitierpflicht – keine Antwort ohne Quellenangabe auf den Datenpfad; „weiß ich nicht“ ist zulässig
M4	Entscheidungspunkt-Wecker	Frühindikator + Schwelle + vereinbarte Reaktion; KI überwacht und meldet mit Kontext statt nackter Alerts	Alarm-Müdigkeit (SOC-Lektion)	Schwellen und Reaktionen werden im Ritual vereinbart, nicht von der KI erfunden – ein Ulysses-Vertrag (Duke): Das Board bindet vorab sein späteres Ich im Affekt
M5	Adressatengerechte Zustellung	Rolle, Kanal, Zeitpunkt, Sprache ; inklusive Lagebild-Abo : periodische rollenspezifische Digests als planbare Push-Form	Bringschuld-Illusion: „zugestellt“ ≠ „verstanden“	Backbrief-Stichproben im Ritual
M6	Lagevortrag-Entwurf	KI entwirft den BLUF-Vortrag für den Sync; der Mensch prüft, kürzt, trägt vor	Das Ritual verkommt zur Verlesung von KI-Text	Der Vortragende zeichnet verantwortlich; Entwurf ≠ Vortrag
M7	Institutionelles Gedächtnis	AAR-Archiv, Decision-/Assumption-Log durchsuchbar und befragbar – „Hatten wir das schon einmal? Was haben wir damals beschlossen – und warum?“	Veraltete Beschlüsse als aktuelle Wahrheit	Jede Antwort trägt Datum und Gültigkeitsstatus der Quelle

Dauerverfügbarkeit vs. Ritual. Entwertet der 24/7-Lagedienst den Portfolio-Sync? Nein – er wertet ihn auf: Wenn die Information vorher bei allen war, wird der Sync vom Berichts- zum **Entscheidungsstermin** (McChrystals O&I war gerade *täglich*). Gefährlich ist nur, das Ritual zu streichen, „weil ja alles im Chat steht“ – dann stirbt die Bühne, auf der Führung Kultur vorlebt (ToT-Erfolgsfaktor 2: Rituale vor Tools).

Ein Bild, n Lesarten. Wie viel Personalisierung verträgt die eine Wahrheit? Lesarten dürfen sich in Auswahl, Tiefe und Sprache unterscheiden – **nie in den Zahlen und nie im Risikobild**. Der gemeinsame Kern (Ampeln, Top-Risiken, Konfidenz) ist unpersonalisierbar; wer ihn personalisiert, baut Wassermelonen-Reports mit KI-Beschleunigung.

Die dritte Spannungsfrage – radikale Transparenz vs. Schutz der Beschäftigten – ist so grundsätzlich, dass sie nicht als Abwägung, sondern als Prinzip beantwortet wird: siehe P6.

C.4 Die sieben KI-Lagebild-Prinzipien

Anschlussfähig an die fünf Lagebild-Prinzipien der EA-Denkschrift – erweitert um ein sechstes (Aggregat-Regel, seit v1.0) und in v2.0 um ein siebtes (Entscheidungshygiene, aus dem Urteilkraft-Block):

1 KI ist Stabsmitglied, nicht Kommandeur.

Sie sammelt, entwirft, übersetzt und widerspricht – sie entscheidet nicht und verantwortet nichts. Verantwortung ist nicht delegierbar an etwas, das keine Rechenschaft tragen kann. (ToT B.1; NATO-Prinzipien; Prediction Machines.)

2 Herkunft vor Eloquenz.

Jede KI-Aussage im Lagebild trägt Quelle und Konfidenz – sonst ist sie die perfekte Information Masquerade. (Milchman; EA-Prinzip 4; DORA-Vertrauenslücke; Art. 50 Abs. 4 KI-VO – menschliche Prüfung und redaktionelle Verantwortung sind zugleich die Voraussetzung, unter der KI-unterstützte Texte keiner Kennzeichnung bedürfen.)

3 Das LLM textet, es rechnet nicht.

Zahlen entstehen im deterministischen Kern (Pipeline, Simulation); die KI formuliert, übersetzt und erklärt sie. Eine frei erfundene Zahl ist ein Systemfehler, kein Schönheitsfehler. (Eigene Architektur; Hubbard-Kalibrierung.)

4 Bereitstellen heißt priorisieren.

Zugestellt wird, was Entscheidungen ändern kann (EVI) – sonst automatisiert KI nur den KPI-Friedhof und produziert Alarm-Müdigkeit. (Hubbard; SOC-Lektion; Goodhart-Warnung.)

5 Der Mensch übt weiter.

Wer jeden Entwurf der KI überlässt, verliert die Kompetenz, ihre Fehler zu erkennen – Ausbildung, Katas und KI-freie Übungen sind der Deskillung-Schutz. (Bainbridge; Marquet „Competence“; Martin „Übung“.)

6 Das Lagebild kennt Teams, keine Personen.

Aggregation auf Team-/ART-Ebene ist harte Regel, kein Stilmittel: kein Personen-Drilldown, keine Individualmessung – erst recht nicht, wenn KI das Auswerten billig macht. Zugleich die wirksamste Compliance-Vereinfachung (Mitbestimmung, Datenschutz, KI-Verordnung: Wer keine Personen bewertet, bleibt außerhalb des Hochrisiko-Bereichs von Anhang III Nr. 4). (Linie der Reihe; Rechts-Prüfbox D.2.)

7 Erst unabhängig urteilen, dann die KI fragen – und immer beide Fehler prüfen. (neu in v2.0)

Urteile leiden an Verzerrung **und** Streuung; KI dämpft beides nur, wenn sie nicht zum Ersatz des eigenen Urteils wird. Deshalb gilt Entscheidungshygiene: unabhängige Einschätzung *vor* dem KI-Vorschlag, Aggregation mehrerer Urteile, strukturierte Teilbewertung, relative statt absolute Skalen – und Noise-Audits, um zu wissen, wie laut die eigene Organisation ist. Ein KI-Vorschlag, der die Streuung der Menschen ersetzt, ohne seine eigene auszuweisen, hat Noise nicht reduziert, sondern versteckt. (Kahneman/Sibony/Sunstein; Duke Truthseeking; Position des Kurators, A2.2.)

C.5 Kompetenzfeld 7: KI-gestützte Stabsarbeit

Die ToT-Denkschrift definiert sechs Kompetenzfelder für Portfolio-Stäbe. KI fügt ein siebtes hinzu – **nicht** „Prompt-Tricks“, sondern Stabshandwerk (und zugleich das, was Art. 4 der KI-Verordnung seit Februar 2025 von jedem Betreiber verlangt: die KI-Kompetenz der eigenen Leute zu fördern):

- **Beauftragen können:** präzise Aufträge an KI formulieren (Kontext, Quellen, Format, Grenzen) – die zivile Form des Befehlsentwurfs.
- **Verifizieren können:** Stichproben ziehen, Quellenangaben prüfen, Zahlen gegen den deterministischen Kern halten; wissen, *wann* volle Prüfung nötig ist (Entscheidungsvorlagen) und wann Stichprobe genügt (Routinetexte).
- **Konfidenz lesen und schreiben:** kalibrierte Unsicherheitssprache verstehen und von KI einfordern (Anschluss an das Kalibrierungstraining aus Curriculum-Modul 3); Prognosen ohne „Resulting“ bewerten (Duke): Ein eingetretenes 15-%-Szenario widerlegt keine P85-Prognose.
- **Automation Bias erkennen:** die eigene Neigung kennen, flüssigen Vorschlägen unter Zeitdruck zu folgen – und Gegenrituale beherrschen (Vier-Augen-Prinzip, anonyme Vorab-Voten *vor* dem KI-Vorschlag, Truthseeking-Gegenprobe in der Gruppe nach Duke).
- **Die gezackte Grenze kartieren:** durch begleitete Praxis lernen, wo KI im eigenen Kontext stark und wo sie unzuverlässig ist (Mollick) – dieses Wissen veraltet mit jedem Modellwechsel und muss gepflegt werden. *Position des Kurators (v2.0):* KI deshalb **stark einbinden**, um die Grenzen überhaupt zu sehen – aber nie ohne Mensch in der Schleife; die Modelle entwickeln sich zu schnell, um zu gamblen.
- **Entscheidungshygiene praktizieren (neu in v2.0):** die eigene Einschätzung notieren, *bevor* man die KI fragt; bei wichtigen Urteilen mehrere unabhängige Einschätzungen aggregieren; Bewertungen in Teilurteile zerlegen; wissen, dass Bias und Noise zwei verschiedene Fehler sind – und dass ein flüssiger KI-Vorschlag beide verdecken kann.
- **Die Zerlegung beherrschen (neu in v2.0):** eine anstehende Entscheidung in Vorhersage, Urteil, Aktion, Ergebnis, Input-Daten und Feedback zerlegen können (Prediction Machines) – und je Bestandteil sagen, was Software, KI und Mensch übernehmen (C.6).

C.6 Wo Software, KI und Mensch je Entscheidungsbestandteil hingehören (neu in v2.0)

Die Zerlegung aus A2.1 macht die Kernfrage dieser Schrift präzise beantwortbar. Für eine typische Portfolio-Entscheidung – „Ziehen wir Vorhaben X vor, obwohl Y dann rutscht?“ – ergibt sich, je Bestandteil:

Bestandteil	Was unsere Software heute leistet	Wo KI helfen kann (Stufe 3, mit Guardrails)	Was beim Menschen bleibt
Input-Daten	Extraktion, Transformation, Kanonisierung, Konfidenz-Flags (Pipeline; Phasen A–C der Roadmap)	Klassifikation, Dubletten, Lückenfindung, Begründung von Qualitäts-Flags (Rolle 1)	Entscheiden, <i>welche</i> Daten zählen – der Informationsbedarf (EVI, Steuern-Phase)
Vorhersage	Monte-Carlo-Forecasts mit Konfidenzintervallen (Simulate), Trend-/Anomalie-Statistik – deterministisch, reproduzierbar	Szenario-Varianten formulieren, Forecast-Ergebnisse <i>erklären</i> , Annahmen hinterfragen (Rolle 3) – nie die Zahl selbst erzeugen (Prinzip 3)	Kalibrierung prüfen; wissen, dass ein Forecast eine Verteilung ist, kein Versprechen (Resulting-Regel)
Urteil	Kennzahlen-Set (Flow + Outcome + Qualität + Mensch) als Bewertungsgrundlage; ROAM/Decision-Log als Struktur	Bewertungskriterien vollständig auflisten, Trade-offs gegenüberstellen, Gegenargumente liefern (Red-Team-Assistent D5)	Das Urteil selbst: was ist Treffer, was Fehler wert – Prinzip 1 und 7; unabhängig urteilen, bevor die KI spricht
Aktion	Kein Automatismus – bewusst	Entscheidungsvorlage und Delta-Briefing formulieren (Rolle 2, D1/D2)	Entscheiden und verantworten; Ausführung nur mit Freigabe (Regler 4)

Ergebnis	Report-Diff zwischen zwei Ständen, Zielerreichung je Kennzahl (D2)	Abweichungen erklären, AAR-Fragen vorbereiten (Rolle 2/3)	AAR mit Resulting-Schutzregel: Entscheidungs- ≠ Ergebnisqualität
Training/Feedback	Historische Reports, Decision-/Assumption-Log, AAR-Archiv als Datenbasis (M7)	Muster über Vorhaben hinweg erkennen, „Hatten wir das schon?“ beantworten (M7, mit Zitierpflicht)	Lernen als Organisation: Kalibrierung nachhalten, Annahmen mit Verfallsdatum prüfen

Drei Lehren aus der Tabelle: **(1)** Die Software steht dort am stärksten, wo Determinismus zählt – Daten, Vorhersage, Ergebnismessung; genau dort darf KI *erklären*, aber nicht *ersetzen*. **(2)** KI wirkt am stärksten in den Übergängen – Rohmaterial für Urteil, Sprache für Aktion, Muster für Feedback – und überall mit Herkunftspflicht. **(3)** Das Urteil ist die einzige Zeile, in der die Spalte „Mensch“ fett steht: Es ist der Bestandteil, dessen Wert mit jeder billigeren Vorhersage steigt – und der Bestandteil, in dem Noise am meisten schadet. Wer die Zerlegung ernst nimmt, weiß deshalb auch, was die Roadmap **nicht** bauen wird: einen Empfehlungs-Button, der die Urteils-Zeile automatisiert.

Teil D – Umsetzung

D.1 Einsatzfelder nach Reifegrad (Entscheider-Matrix)

Einsatzfeld (Rolle)	Reifegrad	Nutzen	Haupt-Guardrail
Datenextraktion & Klassifikation (Rolle 1)	produktiv	Erhebungsaufwand sinkt drastisch	Qualitäts-Flags statt blindem Vertrauen
Zusammenfassung, Lesarten, Übersetzung (Rolle 2)	produktiv mit Prüfpflicht	Rosetta-Stone-Arbeit wird bezahlbar	Prinzip 2 + 3; Prüfpflicht vor Versand
Anomalie/Trend/Forecast (Rolle 3, Stufe 2)	produktiv (Statistik)	Frühindikatoren statt Rückspiegel	Konfidenzintervalle statt Punktwerte
Szenario-/Premortem-Assistenz (Rolle 3, Stufe 3)	erprobenswert	Red-Team-Kapazität für jeden Stab	Mensch moderiert; KI liefert Rohmaterial
Befragbares Lagebild, Wecker, Zustellung (Rolle 4)	erprobenswert	Vom Bericht zum Lagedienst	Zitierpflicht; EVI-Priorisierung; P6
Agentische Erhebung/Aktion (Stufe 4)	Ausblick	Perspektivisch: stehende Recherche-Aufträge	Nur mit Ausführungs-Freigaben (C.2, Regler 4)

D.2 Risiken & Governance

Risiko	Mechanismus	Gegenmittel
Halluzination	LLM erfindet plausible Details, wo Daten fehlen	Zitierpflicht (M3), Generierung nur aus übergebenen Daten, „weiß nicht“ belohnen, Prinzip 3
Automation Bias	KI-Vorschläge werden unter Zeitdruck unkritisch übernommen	Konfidenzausweis, Stichproben-Routine, Vier-Augen bei Entscheidungsvorlagen, Truthseeking-Gegenprobe (Duke)
Noise-Verstärkung (neu in v2.0)	Nicht-deterministische Modelle streuen selbst (Formulierungs-Sensitivität, Halluzination, falsche Quellen); blindes Vertrauen importiert diese Streuung als vermeintliche Objektivität – Bias und Noise werden zugleich unsichtbar	Prinzip 7: unabhängiges Urteil <i>vor</i> der KI, Aggregation, strukturierte Teilbewertung; Noise-Audit (dieselbe Frage mehrfach/mehreren Beurteilern stellen, Streuung messen – auch die der KI); Zitierpflicht gegen falsche Quellen
Deskilling	Automatisierung entzieht die Übung, die Ausnahmen erfordern (Bainbridge)	KI-freie Katas (D.4), Rotation „einmal selbst rechnen“, Prinzip 5
Scheinsicherheit	Perfekt formulierter Text suggeriert Gewissheit im Clausewitz'schen Nebel	Konfidenz gehört zur Zahl <i>und</i> zur Formulierung; normierte Unsicherheitssprache
Goodhart-Verstärkung	Teams optimieren auf KI-lesbare Signale, wenn KI Reports liest und schreibt	Metrik-Set statt Einzelmetrik; qualitative Gegenproben im Ritual
Vertraulichkeit & Datenabfluss	Das Lagebild bündelt sensible Unternehmensinformation; externe KI-Dienste als Abflusskanal	Explizite Deployment-Entscheidung (on-prem / EU-Cloud / Zero-Retention); Datenklassifizierung vor KI-Anbindung

Prompt Injection	Manipulierte Inhalte in Quelldaten steuern die Antworten des befragbaren Lagebilds	Eingaben als Daten behandeln, Rechtstrennung, Antwort-Provenienz
Modell-Drift & Betrieb	KI-Bausteine haben einen eigenen Lifecycle (MLOps)	Betriebsverantwortung benennen; Monitoring der KI-Komponenten selbst
Regulatorische Fehleinstufung (neu in v2.1)	Ein Lagebild, das KI-Auswertungen auf Personenebene zulässt, rutscht in den Hochrisiko-Bereich der KI-Verordnung (Anhang III Nr. 4: Leistung/Verhalten von Beschäftigten) – mit Konformitätsbewertung, Registrierung und Aufsicht; fehlende Transparenzhinweise verletzen Art. 50	Prinzip 6 als Architekturregel im deterministischen Kern (Aggregat-Grenze, keine Personen-Drilldowns – nicht erst in der KI-Schicht); KI-Hinweise ab Tag 1 (Art. 50); Rechts-Prüfbox unten

WARNBOX – DIE MASQUERADE-FALLE

Ein KI-generiertes Lagebild ohne Herkunftsnachweis ist die perfekte Information Masquerade: flüssig, überzeugend, unprüfbar. Es fühlt sich wie ein Informationsgewinn an und ist ein Kontrollverlust. Wer nur eine einzige Regel aus diesem Dokument übernimmt, nehme die **Zitierpflicht**.

Rechts-Prüfbox (bei Einführung juristisch prüfen; Stand der Angaben: 18.08.2026): KI-Verordnung (EU) 2024/1689 („AI Act“) – gestufter Geltungsbeginn: Verbote (u. a. Emotionserkennung am Arbeitsplatz) und die **KI-Kompetenzpflicht (Art. 4)** seit 02.02.2025, Pflichten für Allzweck-KI-Modelle seit 02.08.2025, **Transparenzpflichten (Art. 50) seit 02.08.2026** (eine Übergangsfrist bis 02.12.2026 gilt nur für die maschinenlesbare Markierung durch *Anbieter* nach Abs. 2 bei Bestandssystemen); die Hochrisiko-Pflichten wurden mit dem *Digital Omnibus* (2026) verschoben – Anhang III auf den 02.12.2027, Anhang I auf den 02.08.2028. Vier Punkte betreffen den KI-gestützten Lagedienst unmittelbar:

(1) Betreiber sind wir. Wer ein KI-System beruflich einsetzt, ist *Betreiber* (Art. 3 Nr. 4) – mit der Pflicht, die KI-Kompetenz der eigenen Leute zu fördern (Art. 4; Kompetenzfeld 7 und Modul 7 sind genau das) und mit den Transparenzpflichten des Art. 50.

(2) Transparenz. Ein befragbares Lagebild muss erkennen lassen, dass eine KI antwortet (Art. 50 Abs. 1). KI-erzeugte Texte, die die Öffentlichkeit informieren, sind zu kennzeichnen – *es sei denn*, sie wurden substantiell menschlich geprüft und eine Person trägt die redaktionelle Verantwortung (Art. 50 Abs. 4): Genau das leisten Zitierpflicht, Vier-Augen-Prinzip und der Transparenzhinweis dieser Reihe. Die Kommissionsleitlinien (Juli 2026) setzen die Latte hoch: cursorische Freigabe genügt nicht, und eine substantielle KI-Bearbeitung *nach* der Prüfung lässt die Ausnahme entfallen – Freigabe ist der letzte Schritt.

(3) Die Hochrisiko-Falle. Systeme mit Bezug zu Leistung/Verhalten von Beschäftigten (Anhang III Nr. 4) tragen die volle Pflichtenlast (Risikomanagement, Konformitätsbewertung, Registrierung, menschliche Aufsicht). Prinzip 6 („Teams, keine Personen“) hält das Lagebild davon fern – als Architekturregel im Kern, nicht als Bitte an die KI – und ist damit die wirksamste Compliance-Vereinfachung überhaupt.

(4) Open Source befreit nicht. Die Ausnahme für quelloffene KI (Art. 2 Abs. 12) gilt ausdrücklich nicht für Art. 50 und nicht für Hochrisiko-Systeme – auch unsere Software baut die Transparenzregeln ein (Merkposten unten).

Dazu national: **§ 87 Abs. 1 Nr. 6 BetrVG** (Mitbestimmung bei technischen Einrichtungen zur Leistungs-/Verhaltenskontrolle) und **DSGVO** – wer keine Personendaten auswertet, hat auch hier die halbe Prüfung hinter sich. Betriebsrat früh einbinden, nicht nachträglich. (Quellen: Verordnung (EU) 2024/1689; Leitlinien der Kommission zu Art. 50 und Verhaltenskodex zur Transparenz KI-generierter Inhalte, Juli 2026 – siehe Quellenverzeichnis.)

D.3 Der 90-Tage-Pilot

1. **Woche 1–2 – Steuern:** Informationsbedarf der Entscheider nach EVI erheben („Welche Information würde welche Entscheidung ändern?“); Datenklassifizierung und Deployment-Entscheidung für KI-Dienste treffen; Betriebsrat informieren (P6 als Zusage).
2. **Woche 3–6 – erstes Muster produktiv:** das **Delta-Briefing (M2)** einführen – deterministischer Report-Diff plus KI-Formulierung, vor jedem Sync an alle Teilnehmer. Geringstes Risiko, sofort spürbar, sauberer Beleg für Prinzip 3.
3. **Woche 7–12 – Neuland kontrolliert:** das **befragbare Lagebild (M3)** als überwachter Pilot für einen definierten Nutzerkreis, mit Zitierpflicht und Protokollierung; parallel Kompetenzfeld-7-Schulung für den Stab. *Neu in v2.0 – der Noise-Audit (Woche 9–10):* Eine reale Priorisierungsfrage wird (a) von fünf Stabsmitgliedern unabhängig beurteilt, (b) der KI fünfmal in leicht variiert formuliert, (c) erst danach in der Gruppe diskutiert. Gemessen wird die Streuung in allen drei Runden. Das Ergebnis ist die ehrlichste Zahl des Piloten: Wie laut ist unser Stab – und wie laut ist unsere KI?
4. **Abschluss – Retro auf den Lagedienst selbst:** Vertrauen die Nutzer den Antworten – und *warum* (nicht)? Welche Fragen konnte das Lagebild nicht beantworten (= Lücken im Datenbestand)? Was hat der Noise-Audit gezeigt – und welche Hygiene-Regel wird ab jetzt Routine? Entscheidung über Ausbau (M4, M5, M7).

D.4 Curriculum-Modul 7: KI-gestützte Stabsarbeit

Erweiterung des 12-Monats-Curriculums der ToT-Denkschrift (dort Teil D.2) um ein siebtes Modul – berufsbegleitend, übungslastig:

- **Inhalte:** die vier Rollen und der Automatisierungsregler; Beauftragen/Verifizieren/Konfidenz (C.5); Risikokatalog und Rechtslage (D.2); die gezackte Grenze im eigenen Kontext; *neu in v2.0:* die Zerlegung einer Entscheidung in sechs Bestandteile (A2.1, C.6) und Entscheidungshygiene (A2.2, Prinzip 7).
- **Übungen:** (1) denselben Lagevortrag einmal mit, einmal ohne KI erstellen und vergleichen; (2) ein präpariertes KI-Briefing mit eingebauten Fehlern auditieren („Finde die Halluzination“); (3) Premortem mit KI-Rohmaterial moderieren; (4) **KI-freie Kata** monatlich – eine Lageanalyse in 30 Minuten von Hand, damit die Form sitzt, wenn das Werkzeug ausfällt (Bainbridge-Antwort); *neu in v2.0:* (5) **Zerlegungs-Kata** – eine anstehende Entscheidung in 20 Minuten in die sechs Bestandteile zerlegen und je Bestandteil Software/KI/Mensch zuordnen; (6) **Mini-Noise-Audit** – dieselbe Bewertungsfrage erst still und einzeln, dann mit KI, dann in der Gruppe; Streuung sichtbar machen und besprechen.
- **Artefakt:** eine einseitige **KI-Nutzungsordnung für Stabsarbeit** (welche Stufe wofür, Prüfpflichten, Zitierpflicht, P6, *neu:* Hygiene-Regel „eigenes Urteil zuerst“) – das Governance-Dokument, das die meisten Organisationen noch nicht haben.

D.5 Der einkaufbare Markt (herstellernerneutral, nach Kategorien)

- **Enterprise-LLM-Plattformen** mit Unternehmens-Datenschutz (private Instanzen, EU-Optionen, Zero-Retention-Verträge) – die Basis für Rolle 2 und 4.
- **BI-/Analytics-Copiloten** (natürlichsprachliche Abfrage vorhandener Dashboards) – der schnellste Weg zu M3-ähnlicher Funktionalität, aber Zitierpflicht und Datenpfad prüfen.
- **AIOps-/Anomalie-Werkzeuge** (Betriebsdaten, Stufe 2) – reif; liefern Frühindikatoren für das technische Gesundheitsbild (Anschluss an die SLO/DORA-Dimension der EA-Roadmap).
- **Forecasting-Werkzeuge** (Monte-Carlo/probabilistisch) – vorhanden als eigenes Simulate-Modul; extern nur nötig, wenn Fremd-Datenquellen dazukommen.
- **KI-Kompetenz-Trainings** – auf Stabsarbeit gemünzt einkaufen (Verifikation, Bias, Recht), nicht als Prompt-Kurs; Methodentraining (Red Teaming, Kalibrierung, Superforecasting) siehe Programmtabelle der ToT-Denkschrift D.1.

Anschlussfähigkeit an unsere Software (Merkposten)

Die Roadmap des `portfolio`-Moduls ([docs/portfolio_roadmap_solution-lagebild.md](#)) wird um **Phase D – KI-Assistenz** ergänzt, konsequent **hinter** dem deterministischen Kern:

#	Feature	Muster	Voraussetzung
D1	LLM-Executive-Summary (Zahlen aus <code>summary.py</code> , LLM formuliert)	M1/M6	–
D2	Delta-Briefing (deterministischer Report-Diff + Narration) ← <i>erster Pilot</i>	M2	zwei Report-Stände
D3	Befragbares Lagebild (Q&A über Report-Daten mit Zitierpflicht)	M3	Provenienz-Modell
D4	Anomalie-Hinweise (Statistik im Kern, KI erklärt)	M4	–
D5	Red-Team-Assistent (Premortem-Fragen aus Annahmen-/Decision-Log)	Rolle 3	Roadmap B4
D6	Mehrsprachige Report-Ausleitung	M5	–
D7	Noise-Audit-Unterstützung (n-fach variierte Anfrage, Streuungsausweis) (<i>neu in v2.0</i>)	Prinzip 7	D3

Architektur-Grundsatz: Die KI-Schicht ist optional und austauschbar (eigenes Modell/eigener Schlüssel, lokal oder API); der Kern bleibt offline lauffähig und stdlib-treu. **Kein Kernfeature darf KI voraussetzen – KI veredelt, sie trägt nicht.**

Konsequenzen aus dem Urteilskraft-Block für Phase D (*neu in v2.0; in die Roadmap übernommen*): **(a)** Jede D-Funktion wird einem Entscheidungsbestandteil aus C.6 zugeordnet – D1/D2/D6 der *Aktion* (Formulieren), D3/D7 der *Vorhersage/dem Feedback* (Erklären), D4 dem *Ergebnis*, D5 dem *Urteil* (Rohmaterial, nie das Urteil selbst). **(b)** Das befragbare Lagebild (D3) bekommt eine **Hygiene-Vorgabe**: Es kann konfiguriert werden, die eigene Einschätzung des Fragenden *vor* der Antwort abzufragen und beide nebeneinanderzustellen. **(c)** Neu **D7 · Noise-Audit-Unterstützung**: dieselbe Frage n-fach variiert stellen, Antwort-Streuung protokollieren und ausweisen – die KI misst ihre eigene Lautstärke. **(d)** Bestätigt: kein Empfehlungs-Button für die Urteils-Zeile.

Rechtliche Leitplanken für Phase D (*neu in v2.1; aus der Rechts-Prüfbox D.2, in die Roadmap übernommen*): **(a)** Das befragbare Lagebild (D3) zeigt an, dass eine KI antwortet (Art. 50 Abs. 1 KI-VO); KI-formulierte Report-Texte (D1/D2/D6) tragen einen sichtbaren KI-Hinweis, bis ein Mensch sie geprüft und freigegeben hat – die Freigabe ist zugleich die Ausnahme von der Kennzeichnungspflicht (Art. 50 Abs. 4). **(b)** Die Aggregat-Grenze (Prinzip 6) wird im deterministischen Kern erzwungen: Kein D-Feature erhält Personendaten – so bleibt das Lagebild außerhalb von Anhang III Nr. 4. **(c)** Die Open-Source-Lizenz befreit nicht von Art. 50 (Art. 2 Abs. 12) – die Hinweise sind Teil des Produkts, nicht der Konfiguration. **(d)** Betreiber-Nachweis: Modell, Version, Systemvorgaben und Datenklassifizierung je Deployment werden protokolliert – die Grundlage, um Kompetenz- und Transparenzpflicht belegen zu können.

Reflexiver Beleg: Diese Reihe praktiziert, was sie beschreibt. Jede Denkschrift entsteht KI-gestützt und trägt den Transparenzhinweis (Versionierungs-Grundsätze 5 und 6 der Reihe: KI-Unterstützung benannt, Inhalt geprüft und redaktionell verantwortet, Freigabe des Kurators als letzter Schritt vor der Veröffentlichung – zugleich die Voraussetzungen der Ausnahme in Art. 50 Abs. 4 KI-VO); die [Handbücher der Software](#) erscheinen KI-übersetzt in fünf Sprachen; die Forecasts rechnet ein deterministisches Monte-Carlo-Modul, dessen Ergebnisse KI erklären, aber nie erzeugen darf. Wir schreiben nicht über KI-Stabsarbeit – wir betreiben sie, mit genau den Leitplanken, die dieses Dokument fordert.

Anschluss an die Reihe

- **ToT-Denkschrift:** Das O&I-Ritual bleibt das Entscheidungsforum – der Lagedienst versorgt es; Empowered Execution wird gegen den Rezentralisierungs-Reflex verteidigt (B.3/RAND); Kompetenzfeld 7 und Modul 7 docken an Curriculum und Kompetenzkatalog an.
- **EA-Denkschrift:** Die sieben KI-Prinzipien setzen die fünf Lagebild-Prinzipien fort; die EA-Dimensionen (Capability, NFR, ROAM, Dependencies) sind die Datenbasis, die der Lagedienst befragbar macht; der Regelkreis bekommt mit M2/M4 seine Taktgeber.
- **Nebenschrift Architektur-Vordenker:** Hohpes „first derivative“ wird als Delta-Briefing zum Produkt; Fowlers Muster-Vokabular begründet die gemeinsame Sprache der KI-Nutzungsordnung; Martins Übungskultur trägt Prinzip 5; Hohmanns Rituale bleiben die menschliche Bühne, auf der KI-Rohmaterial verhandelt wird.
- **Lagebild-Literaturliste** (im Werkstatt-Archiv des Kurators): nimmt die neuen Quellen zu KI & Lagedienst und zur Urteilkraft (Duke) auf.

Quellenverzeichnis

Eigener Korpus (Volltexte im Werkstatt-Archiv des Kurators): *Architecting Enterprise AI Applications · Mastering Machine Learning Architecture and Solutions · AI and Microservices · Fault Detection in Microservice Architectures* (alle Apress) · Milchman: *Unit Oriented Enterprise Architecture – Constructing Large Sociotechnical Systems in the Age of AI*. Apress, 2024.

Neu archiviert (frei verfügbar, PDFs im Werkstatt-Archiv des Kurators):

- NATO ACT: *NATO Decision-Making in the Age of Big Data and Artificial Intelligence* (2021) – act.nato.int
- RAND: *How Artificial Intelligence Could Reshape Four Essential Competitions in Future Warfare* (RRA4316-1) – rand.org
- Bainbridge, Lianne: „Ironies of Automation“. *Automatica* 19(6), 1983 – doi.org.

Online (zitiert):

- DORA: *State of AI-assisted Software Development* (2025) – dora.dev/dora-report-2025 (Download über Formular)
- US Joint Chiefs of Staff: *Joint Publication 2-0, Joint Intelligence* – jcs.mil/Doctrine (öffentliche Fassung)
- Atlantic Council: *How NATO can integrate AI to prevail in future algorithmic warfare* – atlanticcouncil.org
- NATO: *Summary of the NATO Artificial Intelligence Strategy / Principles of Responsible Use* (2021) – nato.int
- Parasuraman/Sheridan/Wickens: „A Model for Types and Levels of Human Interaction with Automation“. *IEEE Trans. SMC-A* 30(3), 2000 – doi.org.

Ausgewertet ergänzend (neu in v1.1): Duke, Annie: *Thinking in Bets. Making Smarter Decisions When You Don't Have All the Facts*. Portfolio, 2018 — Teil A2.7; ausführliche Würdigung: ToT-Denkschrift (A2.5).

Der Urteilkraft-Block (gelesen und bewertet durch den Kurator, v2.0): Agrawal, Ajay / Gans, Joshua / Goldfarb, Avi: *Prediction Machines – The Simple Economics of Artificial Intelligence*. Harvard Business Review Press, 2018; Updated and Expanded 2022. · Kahneman, Daniel / Sibony, Olivier / Sunstein, Cass R.: *Noise – A Flaw in Human Judgment*. Little, Brown Spark, 2021 (dt. *Noise – Was unsere Entscheidungen verzerrt*, Siedler). · Mollick, Ethan: *Co-Intelligence – Living and Working with AI*. Portfolio, 2024 (dt. *Co-Intelligenz*, Redline 2025).

Rechtsrahmen (Stand 18.08.2026; Angaben ohne Gewähr, bei Einführung juristisch prüfen): Verordnung (EU) 2024/1689 (KI-Verordnung / AI Act) · eur-lex.europa.eu/eli/reg/2024/1689/oj · Europäische Kommission: Leitlinien zu den Transparenzpflichten nach Art. 50 (Juli 2026) und *Code of Practice on Transparency of AI-Generated Content* · digital-strategy.ec.europa.eu · Digital-Omnibus-Änderung der KI-Verordnung (2026; verschobene Hochrisiko-Fristen, Übergangsfrist für Art. 50 Abs. 2) · national: BetrVG, DSGVO.

Interne Bezüge: „Team of Teams & Stabsausbildung“ · „Enterprise-Architektur & Solution-Lagebild“ · „Architektur-Vordenker“ · Beschlussgrundlage im Werkstatt-Archiv des Kurators (Beschlussstand 13.08.2026) · Lagebild-Literaturliste (im Werkstatt-Archiv des Kurators) · [docs/portfolio_roadmap_solution-lagebild.md](#) (Phase D).

Versionshistorie

Version	Datum	Änderungen
1.0	13.08.2026	Erstfassung gemäß Beschlussstand der Diskussionsgrundlage vom 13.08.2026: Doppel-Leitbild (Stabsmitglied/Lagedienst), Vier-Stufen-Taxonomie, Teil B Nachrichtenwesen + NATO/RAND + SOC, sieben Bereitstellungs-Muster, sechs KI-Prinzipien (inkl. P6 „Teams, keine Personen“), Automatisierungsregler, Kompetenzfeld 7, 90-Tage-Pilot, Curriculum-Modul 7, Software-Phase D (Pilot: Delta-Briefing). Geplant v2.0: Vertiefung nach Auswertung von <i>Prediction Machines</i> , <i>Co-Intelligence</i> , <i>Noise</i> .
1.1	13.08.2026	Kleinere Ergänzung (keine Neuausrichtung): Annie Duke, <i>Thinking in Bets</i> , als A2.6 eingearbeitet – Resulting-Lesekompetenz für probabilistische Lageinformation (C.5), Truthseeking-Gegenprobe gegen Automation Bias (C.5, D.2), Entscheidungspunkt-Wecker (M4) als Ulysses-Vertrag eingeordnet (C.3); Schwesterdokument-Verweis auf ToT v4.0 aktualisiert. Die Vertiefung v2.0 bleibt vorgesehen.
1.2	14.08.2026	Verweispflege für die Online-Veröffentlichung: Querverweise auf die Reihe verlinken die Online-Übersicht, frei verfügbare Quellen mit echten URLs (DOI/offizielle Seiten), Werkstatt-Pfade entfernt; Querverweise ohne Versionsnummern. Inhalt unverändert.
2.0	15.08.2026	Vertiefung nach Lektüre von <i>Prediction Machines</i> , <i>Noise</i> und <i>Co-Intelligence</i> durch den Kurator: Teil A2 als Urteilkraft-Block neu gefasst (A2.1 Zerlegung in sechs Bestandteile/AI Canvas; A2.2 Bias ≠ Noise + Entscheidungshygiene; A2.3 gezackte Grenze + Mensch in der Schleife – jeweils mit Position des Kurators; Duke als A2.7 integriert). Neu: Prinzip 7 Entscheidungshygiene (C.4), Kapitel C.6 Zuordnung Software/KI/Mensch je Entscheidungsbestandteil , Kompetenzfeld 7 um Hygiene und Zerlegung erweitert (C.5), Risiko „Noise-Verstärkung“ (D.2), Noise-Audit im 90-Tage-Piloten (D.3), Zerlegungs-Kata und Mini-Noise-Audit in Modul 7 (D.4), Phase D um D7 Noise-Audit-Unterstützung und Hygiene-Vorgabe für D3 erweitert (Roadmap-PR). Executive Summary entsprechend fortgeschrieben.
2.1	18.08.2026	Kleinere Ergänzung (keine Neuausrichtung): Rechtsrahmen der KI-Verordnung auf den Stand August 2026 gebracht – Rechts-Prüfbox in D.2 neu gefasst (Geltungsbeginn nach dem Digital Omnibus, Betreiberpflichten Art. 4/Art. 50 samt Ausnahme bei redaktioneller Verantwortung, Hochrisiko-Falle Anhang III Nr. 4, Open-Source-Ausnahme gilt nicht für Art. 50), Risiko „Regulatorische Fehleinstufung“ (D.2), Bezüge in B.3, Prinzip 2/6 (C.4) und Kompetenzfeld 7 (C.5), rechtliche Leitplanken für Phase D (Merkposten; in die Roadmap übernommen), Quellenverzeichnis „Rechtsrahmen“. Dokumentiert zugleich die am 18.08.2026 noch innerhalb der v2.0 vorgenommene Klarstellung des B.3-Satzes zur NATO-Grenzziehung („KI dient im Stab, sie führt nicht“). Transparenzhinweis nach dem Reichen-Standard (Versionierungs-Grundsätze 5/6).