

KI-Governance im Mittelstand: Rollen, Risiken und Implementierungspfade

Lucien André Reuter

EONAR GmbH | lucien.reuter@eonar.de | <https://lucienreuter.de>

Working Paper

LR-WP-2026-01

29. Januar 2026

Version: v1.0, Stand: 29. Januar 2026

Empfohlene Zitierweise: Reuter, L. A. (2026). *KI-Governance im Mittelstand: Rollen, Risiken und Implementierungspfade*. Working Paper, LR-WP-2026-01. <https://lucienreuter.de/publikationen/ki-governance-mittelstand/>

Dokumenttyp: Working Paper

Reportnummer: LR-WP-2026-01

URL: <https://lucienreuter.de/publikationen/ki-governance-mittelstand/>

PDF: <https://lucienreuter.de/publikationen/ki-governance-mittelstand/ki-governance-mittelstand.pdf>

Zusammenfassung

Sobald KI im Mittelstand über einzelne Pilotprojekte hinauswächst, stellen sich Fragen, die kein Werkzeug beantwortet: Wer gibt einen Anwendungsfall frei, wer verantwortet die zugrunde liegenden Daten, und wie wird menschliche Aufsicht praktisch sichergestellt? Der Beitrag übersetzt dafür die risikobasierte Logik des EU AI Act, die Funktionen des NIST AI Risk Management Framework, die Anforderungen aus ISO/IEC 42001 und die Werteprinzipien der OECD in eine Rollen- und Umsetzungslogik, die auch ohne eigene Governance-Abteilung tragfähig bleibt. Der Rahmen gliedert sich in sechs Dimensionen: Use-Case-Inventar, Risikoklassifizierung, Daten- und Modellverantwortung, Rollen und Entscheidungsrechte, Betriebs- und Monitoringprozesse sowie Kompetenz- und Akzeptanzsicherung. Die Kontrolldichte richtet sich nach dem Risiko des einzelnen Anwendungsfalls, nicht nach der Technologie insgesamt. Leitend ist daher nicht die Frage, welches Werkzeug am meisten verspricht, sondern welcher Anwendungsfall verantwortet werden kann. Der Beitrag erhebt keine Primärdaten und ersetzt keine Rechtsberatung.

Keywords: KI-Governance; Mittelstand; Künstliche Intelligenz; AI Act; Datenverantwortung; digitale Transformation

Kernaussagen

- KI-Governance im Mittelstand sollte nicht als Bremse verstanden werden, sondern als Entscheidungsarchitektur für verantwortliche Nutzung.
- Ein Use-Case-Inventar mit Risikoklassifizierung, Datenverantwortung und Rollenmodell ist die Grundlage für steuerbare KI-Einführung.
- Generative KI und KI-Agenten benötigen zusätzliche Kontrollen für Kontext, Datenzugriff, Testing, menschliche Aufsicht und laufendes Monitoring.

Inhaltsverzeichnis

1	Einleitung	4
2	Problemstellung und Praxisfrage	4
3	Begriffe und Abgrenzung	5
4	Stand der Diskussion	6
5	Methodisches Vorgehen	7
6	Ein sechsdimensionaler Bezugsrahmen für KI-Governance	8
6.1	Dimension 1: Use-Case-Inventar	8
6.2	Dimension 2: Risikoklassifizierung	8
6.3	Dimension 3: Daten- und Modellverantwortung	9
6.4	Dimension 4: Rollen und Entscheidungsrechte	9
6.5	Dimension 5: Betrieb, Monitoring und Incident-Prozess	9
6.6	Dimension 6: Kompetenz, Akzeptanz und menschliche Aufsicht	10
7	Implikationen für die Umsetzung	10
8	Diskussion	11
9	Limitationen	11
10	Fazit	12

1 Einleitung

Künstliche Intelligenz ist in mittelständischen Organisationen kein fernes Zukunftsthema mehr. Sprachmodelle, Assistenzsysteme, automatische Klassifikation, Prognosemodelle, Lernplattformen, Wissenssuche und KI-Agenten werden zunehmend in operative Prozesse eingebettet. Die Einstiegshürden sind niedrig: Viele Werkzeuge lassen sich ohne große Infrastruktur testen, Fachbereiche können eigenständig experimentieren, und erste Produktivitätseffekte sind oft schnell sichtbar. Genau darin liegt aber auch ein Governance-Risiko.

Wenn KI-Anwendungen ohne klare Verantwortlichkeiten entstehen, können Schattenprozesse, Datenschutzprobleme, falsche Erwartungen, unklare Haftungsfragen, Qualitätsrisiken und Akzeptanzprobleme folgen. Ein Chatbot, der interne Dokumente zusammenfasst, wirkt harmlos, kann aber vertrauliche Informationen verarbeiten oder falsche Aussagen überzeugend formulieren. Ein KI-Agent, der Daten aus Systemen abrufen und Workflows anstößt, kann Effizienz erzeugen, benötigt aber klare Grenzen, Berechtigungen, Protokollierung und menschliche Eingriffsmöglichkeiten. Ein Lernsystem kann individuelle Förderung unterstützen, wirft aber Fragen nach Transparenz, Bias, Datennutzung und pädagogischer Verantwortung auf.

Der EONAR-Beitrag zum Aufbau von AI Agents beschreibt praxisnah, dass schon am Anfang Zweck, Scope, Modellwahl, Integrationen, Memory, Orchestrierung, User Interface sowie Testing und Evaluation geklärt werden müssen [8]. Der Beitrag zu KI-gestütztem Lernen betont ergänzend, dass der Einsatz von KI nicht nur eine technische Frage ist, sondern ethische, didaktische, organisatorische und datenschutzbezogene Überlegungen erfordert [4]. Beide Perspektiven lassen sich verallgemeinern: KI-Systeme werden erst dann verantwortbar, wenn sie nicht nur gebaut oder gekauft, sondern geführt, überwacht und in Organisationen eingebettet werden.

Dieses Working Paper entwickelt dafür einen mittelstandstauglichen Governance-Rahmen. Es verbindet regulatorische und standardisierungsnahe Quellen mit praktischen Umsetzungsfragen. Der EU AI Act setzt einen risikobasierten Rechtsrahmen für KI in Europa [2]. NIST stellt mit dem AI Risk Management Framework Funktionen zur Strukturierung von KI-Risiken bereit [5]. ISO/IEC 42001 beschreibt Anforderungen an ein KI-Managementsystem [3]. Die OECD AI Principles formulieren Werte wie Menschenrechte, Transparenz, Robustheit und Verantwortlichkeit [7]. Für KMU geht es nicht darum, diese Rahmen vollständig als Großkonzernprogramm zu kopieren, sondern sie in handhabbare Rollen, Routinen und Entscheidungspunkte zu übersetzen.

2 Problemstellung und Praxisfrage

KI-Governance scheitert in der Praxis häufig an zwei Extremen. Das erste Extrem ist der Tool-Aktionismus. Organisationen testen einzelne Anwendungen, ohne zu klären, welche Daten verarbeitet werden, welche Entscheidungen beeinflusst werden, welche Risiken entstehen und wer verantwortlich ist. Das zweite Extrem ist Überregulierung. Aus Sorge vor Risiken werden starre Freigabeprozesse aufgebaut, die Innovation ausbremsen und in der Praxis umgangen werden.

Für den Mittelstand braucht es einen dritten Weg: eine Governance, die Risiken ernst nimmt, aber klein genug bleibt, um tatsächlich genutzt zu werden. Sie muss zwischen einem einfachen Textassistenzsystem, einem internen Wissenssuchsystem, einem KI-gestützten Lernpfad, einem Prognosemodell, einer automatisierten Entscheidungsvorbereitung und einem agentischen Workflow unterscheiden. Nicht jeder Use Case benötigt dieselben Kontrollen. Aber jeder produktive KI-Einsatz benötigt Mindestklarheit über Zweck, Daten, Verantwortlichkeit, Nutzerinformation, Monitoring und Eskalation.

Die leitende Praxisfrage dieses Working Papers lautet daher:

Wie können mittelständische Organisationen KI-Anwendungsfälle so bewerten, freigeben, betreiben und weiterentwickeln, dass Nutzen, Risiko, Datenverantwortung, menschliche Aufsicht und organisatorische Akzeptanz gemeinsam steuerbar werden?

Diese Frage ist bewusst als Gestaltungsfrage formuliert. Das Paper liefert keine juristische Detailauslegung des AI Act und keine Zertifizierungsanleitung für ISO/IEC 42001. Es entwickelt einen Bezugsrahmen, der mittelständischen Unternehmen hilft, regulatorische, technische und organisatorische Anforderungen in praktikable Governance-Routinen zu übersetzen.

3 Begriffe und Abgrenzung

KI-Governance bezeichnet in diesem Beitrag die Gesamtheit von Entscheidungsrechten, Rollen, Prozessen, Kontrollen und Lernroutinen, mit denen eine Organisation den Einsatz von KI-Systemen verantwortet. Governance ist damit mehr als Compliance. Sie umfasst auch strategische Priorisierung, Datenqualität, Risikobewertung, Betrieb, Monitoring, Schulung und Akzeptanz.

Ein KI-System wird hier breit verstanden als maschinenbasiertes System, das aus Eingaben Ausgaben wie Vorhersagen, Inhalte, Empfehlungen oder Entscheidungen erzeugt, die physische oder virtuelle Umgebungen beeinflussen können. Diese Definition orientiert sich an der OECD-Beschreibung und ist anschlussfähig an moderne regulatorische Definitionen [7]. Für die Praxis ist wichtig, dass nicht nur spektakuläre generative Modelle darunterfallen. Auch Klassifikationsmodelle, Empfehlungssysteme, Prognosemodelle und automatisierte Entscheidungsunterstützung können governance-relevant sein.

Generative KI bezeichnet Systeme, die neue Inhalte wie Text, Bilder, Code, Audio oder strukturierte Antworten erzeugen. NIST weist im Generative AI Profile darauf hin, dass generative KI bestehende Risiken verstärken und neue Risiken erzeugen kann, etwa Konfabulationen, Informationsintegrität, Datenschutz, Informationssicherheit, Bias, Urheberrechtsfragen und Human-AI-Konfiguration [6]. Für KMU ist diese Risikoliste besonders relevant, weil generative KI oft schnell und dezentral eingeführt wird.

KI-Agenten sind in diesem Beitrag KI-gestützte Systeme, die Aufgaben nicht nur beantworten, sondern über Tools, Datenquellen, Speicher und Workflows Handlungen vorbereiten oder

auslösen. Der EONAR-Agentenbeitrag beschreibt dafür Bausteine wie Purpose & Scope, System Prompt Design, Modellwahl, Tools & Integrationen, Memory, Orchestration, User Interface sowie Testing und Evaluation [8]. Governance muss bei solchen Systemen stärker auf Berechtigungen, Logging, Grenzen, Rückfallmechanismen und menschliche Kontrolle achten.

Abzugrenzen ist KI-Governance von reiner IT-Sicherheit und reinem Datenschutz. Beide sind wichtige Teilbereiche, reichen aber nicht aus. Ein KI-System kann sicher betrieben werden und dennoch problematische Empfehlungen erzeugen. Es kann datenschutzkonform sein und dennoch fachlich falsche oder diskriminierende Wirkungen haben. KI-Governance verbindet daher technische, rechtliche, fachliche und soziale Perspektiven.

4 Stand der Diskussion

Die Diskussion um KI-Governance hat sich von allgemeinen Ethikprinzipien zu operativen Rahmenwerken verschoben. Die OECD AI Principles wurden 2019 angenommen und 2024 aktualisiert. Sie betonen unter anderem inklusives Wachstum, Menschenrechte und demokratische Werte, Transparenz und Erklärbarkeit, Robustheit, Sicherheit, Schutz sowie Verantwortlichkeit [7]. Für Unternehmen sind diese Prinzipien wichtig, aber allein noch nicht ausreichend. Sie müssen in konkrete Entscheidungen übersetzt werden: Welche KI-Anwendung ist zulässig? Welche Daten dürfen genutzt werden? Wer prüft Bias? Wer entscheidet über Freigabe? Wer überwacht den Betrieb?

Der EU AI Act setzt für den europäischen Kontext einen risikobasierten Regulierungsrahmen. Nach der europäischen Logik werden KI-Systeme unter anderem nach Risikostufen behandelt; bestimmte Praktiken sind verboten, Hochrisiko-Systeme unterliegen besonderen Anforderungen, Transparenzpflichten gelten für bestimmte Anwendungen und viele risikoarme Anwendungen bleiben weitgehend frei [2]. Für KMU bedeutet dies nicht, dass jede KI-Anwendung ein Hochrisiko-System ist. Es bedeutet aber, dass Organisationen zunächst wissen müssen, welche KI-Systeme sie einsetzen und in welchem Kontext diese wirken.

NIST verfolgt mit dem AI RMF eine freiwillige Risikomanagementlogik. Das Framework strukturiert KI-Risikomanagement über die Kernfunktionen Govern, Map, Measure und Manage [5]. Diese vier Funktionen sind für KMU gut übersetzbar: Govern klärt Verantwortlichkeit und Policies, Map klärt Kontext und Wirkung, Measure bewertet Risiken und Leistungsfähigkeit, Manage steuert Maßnahmen und Nachverfolgung. Der besondere Wert liegt darin, dass KI-Risiken nicht nur technisch gemessen, sondern organisatorisch eingebettet werden.

ISO/IEC 42001 ergänzt diese Perspektive als Managementsystemstandard. Die Norm beschreibt Anforderungen an ein Artificial Intelligence Management System und ist für Organisationen relevant, die KI-basierte Produkte oder Services entwickeln, bereitstellen oder nutzen [3]. ISO betont einen strukturierten Umgang mit Risiken und Chancen, verantwortliche Nutzung, Transparenz und Zuverlässigkeit. Für KMU muss daraus nicht sofort eine Zertifizierung folgen. Die Managementsystemlogik hilft aber, KI-Governance als wiederkehrenden Prozess zu verstehen: planen, umsetzen, prüfen, verbessern.

Generative KI verschärft die Governance-Frage. NIST beschreibt im Generative AI Profile unter anderem Risiken wie Konfabulation, Datenschutz, Informationsintegrität, Bias, Human-AI-Konfiguration, Informationssicherheit und Wertschöpfungskettenrisiken [6]. Diese Risiken sind in mittelständischen Organisationen oft weniger abstrakt, als sie klingen. Sie betreffen beispielsweise falsche Zusammenfassungen, unklare Quellen, Weitergabe vertraulicher Daten an externe Systeme, unkontrollierte Prompt-Workflows oder die Integration von Drittanbieter-Komponenten in interne Prozesse.

Die eingangs eingeführten EONAR-Beiträge liefern hierzu die operative Brücke. Der KI-Lernbeitrag betont, dass KI kein Selbstzweck ist und in pädagogischen Kontexten ethische, didaktische und organisatorische Rahmenbedingungen benötigt. Der Agenten-Leitfaden zeigt, dass selbst technisch orientierte KI-Systeme nur dann produktiv werden, wenn Zweck, Kontext, Schnittstellen, Testlogik und Nutzererfahrung sauber geklärt sind. Zusammengenommen stützen diese Quellen die These: KI-Governance ist nicht nachgelagerte Kontrolle, sondern Voraussetzung für produktive Nutzung.

5 Methodisches Vorgehen

Dieses Working Paper ist ein konzeptioneller Beitrag mit praxisorientiertem Bezugsrahmen. Es basiert auf einer kuratierten Auswertung geprüfter EONAR-Fachbeiträge und externer Primär- und Standardisierungsquellen. Es wurden keine Primärdaten erhoben, keine Interviews geführt und keine statistische Analyse vorgenommen.

Die Quellenbasis wurde im Rahmen der Quellenprüfung bewusst bereinigt. Nicht verwendet werden Beiträge, deren Volltext nicht belastbar zugänglich war oder deren URLs auf andere Inhalte führen. Genutzt werden der EONAR-Beitrag zum KI-gestützten Lernen sowie der EONAR-Beitrag zum Aufbau von AI Agents.

Die externen Quellen wurden nach Belastbarkeit ausgewählt. Für regulatorische Aussagen wird auf Regulation (EU) 2024/1689 und die europäische AI-Act-Systematik verwiesen. Für Risikomanagement wird NIST AI RMF 1.0 sowie das NIST Generative AI Profile verwendet. Für Managementsystemlogik wird ISO/IEC 42001 herangezogen. Für normative Werte und Begriffslogik werden die OECD AI Principles genutzt. Aus diesen Quellen wird ein mittelstandstauglicher Rahmen abgeleitet, ohne daraus eine Rechtsberatung oder Zertifizierungszusage zu machen.

Dieses normen- und literaturgestützte Vorgehen ist von einer fallstudienbasierten Theoriebildung nach Eisenhardt zu unterscheiden [1]. Eisenhardts Verfahren beruht auf iterativem Case-Sampling, systematischer Kodierung und Konstantvergleich zwischen mehreren Fällen; der vorliegende Rahmen wird stattdessen aus Regulierungstexten, Standards und geprüften Praxisbeiträgen abgeleitet. Eine fallstudienbasierte Prüfung, etwa anhand mehrerer KI-Governance-Einführungen in mittelständischen Unternehmen, könnte den Rahmen in einem weiteren Schritt empirisch validieren.

6 Ein sechsdimensionaler Bezugsrahmen für KI-Governance

Der vorgeschlagene Bezugsrahmen besteht aus sechs Dimensionen. Sie können als Checkliste für einzelne KI-Anwendungsfälle oder als Minimalstruktur für ein KI-Governance-Programm genutzt werden. Entscheidend ist, dass die Dimensionen zusammen gedacht werden: Ein Use Case ohne Risikobewertung ist blind; eine Risikobewertung ohne Rollen ist wirkungslos; Rollen ohne Monitoring verlieren nach dem Go-live an Bedeutung.

6.1 Dimension 1: Use-Case-Inventar

KI-Governance beginnt mit Sichtbarkeit. Eine Organisation kann nur steuern, was sie kennt. Ein Use-Case-Inventar sollte alle produktiven, geplanten und relevanten experimentellen KI-Anwendungen erfassen. Dazu gehören interne Assistenten, Chatbots, Automatisierungen, KI-gestützte Analysewerkzeuge, Lernsysteme, Empfehlungssysteme, Prognosemodelle, Dokumentenverarbeitung, Coding-Assistenten und KI-Agenten.

Für jeden Use Case sollten Mindestinformationen dokumentiert werden: Zweck, Fachbereich, verantwortliche Person, Nutzergruppe, eingesetztes Modell oder Produkt, Datenarten, Datenquellen, externe Anbieter, Integrationen, Ergebnisart, Entscheidungsrelevanz, Status, geplante Nutzung und Monitoringverantwortung. Dieses Inventar muss nicht komplex sein. Für den Einstieg reicht oft eine Tabelle. Wichtig ist, dass sie gepflegt und als Entscheidungsgrundlage genutzt wird.

6.2 Dimension 2: Risikoklassifizierung

Die zweite Dimension ordnet Use Cases nach Risiko. Dabei sollte eine Organisation nicht nur auf regulatorische Kategorien schauen, sondern auch auf operative und reputationsbezogene Risiken. Relevante Kriterien sind: Werden personenbezogene oder vertrauliche Daten verarbeitet? Beeinflusst die Anwendung Entscheidungen über Menschen? Kann ein falsches Ergebnis finanzielle, rechtliche oder gesundheitliche Folgen haben? Ist das System nach außen sichtbar? Löst es Aktionen aus? Nutzt es autonome oder agentische Funktionen? Kann es diskriminierende oder irreführende Ergebnisse erzeugen?

Eine einfache Mittelstandslogik kann vier interne Risikostufen unterscheiden: niedrig, beobachtungspflichtig, freigabepflichtig und hoch. Niedrige Risiken betreffen etwa interne Textentwürfe ohne sensible Daten. Beobachtungspflichtige Anwendungen nutzen interne Informationen oder haben größere Reichweite. Freigabepflichtige Anwendungen beeinflussen Prozesse, Kundenkommunikation oder Entscheidungen. Hohe Risiken betreffen sensible Daten, rechtlich relevante Entscheidungen, Beschäftigung, Bildung, kritische Infrastruktur oder automatisierte Handlungen mit erheblicher Wirkung. Diese interne Klassifizierung ersetzt keine rechtliche AI-Act-Prüfung, erleichtert aber die Vorbereitung.

6.3 Dimension 3: Daten- und Modellverantwortung

KI-Systeme sind von Daten, Kontext und Modellverhalten abhängig. Datenverantwortung bedeutet deshalb mehr als Datenschutz. Sie umfasst Datenqualität, Zweckbindung, Zugriffsrechte, Aktualität, Herkunft, Löschlogik, Trainings- oder Kontextdaten und die Frage, ob Informationen an externe Anbieter übertragen werden. Gerade bei generativer KI und KI-Agenten ist zu klären, welche Daten in Prompts, Kontextfenster, Vektordatenbanken, Logs oder Memory-Systeme gelangen.

Modellverantwortung umfasst Modellwahl, Versionierung, Leistungsgrenzen, bekannte Risiken, Evaluationslogik und Lieferantenabhängigkeiten. Bei externen Modellen müssen Nutzungsbedingungen, Datenverarbeitung, Verfügbarkeit, Modellupdates und Dokumentation betrachtet werden. Bei internen oder feinjustierten Modellen kommen Trainingsdaten, Evaluation, Drift und technische Betriebsverantwortung hinzu.

6.4 Dimension 4: Rollen und Entscheidungsrechte

KI-Governance braucht ein Rollenmodell, das schlank genug für den Mittelstand ist. Eine mögliche Mindeststruktur unterscheidet Geschäftsverantwortung, fachliche Verantwortung, technische Verantwortung, Daten- oder Datenschutzverantwortung, Risiko- und Compliance-Perspektive sowie Nutzervertretung. Nicht jede Rolle muss eine eigene Stelle sein. Entscheidend ist, dass die Perspektiven abgedeckt und Entscheidungen dokumentiert werden.

Für produktive KI-Anwendungen sollte klar sein, wer den Use Case beantragt, wer fachlich freigibt, wer technische Integration prüft, wer Daten- und Datenschutzfragen bewertet, wer Nutzerinformation und Schulung verantwortet und wer nach dem Go-live überwacht. Bei höheren Risiken empfiehlt sich ein kleines KI-Board oder ein bestehender Digital-/IT-Lenkungskreis mit KI-Kompetenz. Dieser Kreis sollte nicht jede Prompt-Idee prüfen, aber produktive und risikorelevante Anwendungen freigeben.

6.5 Dimension 5: Betrieb, Monitoring und Incident-Prozess

KI-Governance endet nicht mit der Freigabe. Viele Risiken entstehen im Betrieb: Modellverhalten verändert sich, Datenquellen ändern sich, Nutzer umgehen Regeln, ein Anbieter aktualisiert Modelle, Prompts werden angepasst, Integrationen wachsen oder ein System wird in neuen Kontexten genutzt. Deshalb braucht jede relevante KI-Anwendung Monitoring und einen Eskalationsweg.

Monitoring kann je nach Risiko unterschiedlich aussehen: Stichprobenprüfung, Qualitätsmetriken, Nutzerfeedback, Fehlerberichte, Bias-Prüfungen, Auswertungen von Abbruchraten, Protokollierung kritischer Aktionen, Review von Prompt-Änderungen oder regelmäßige Revalidierung. Bei KI-Agenten sind zusätzliche Kontrollen nötig: Tool-Berechtigungen, Aktionslimits, Genehmigungsschritte, Rücknahmeoptionen, Logging und Tests vor Änderungen.

Ein Incident-Prozess sollte festlegen, was bei falschen Ausgaben, Datenschutzverdacht, diskriminierenden Ergebnissen, Sicherheitsereignissen, Fehlentscheidungen oder unerwarteten

Aktionen passiert. Wer wird informiert? Wer stoppt das System? Wer bewertet den Schaden? Wer kommuniziert intern oder extern? Wer dokumentiert Korrekturmaßnahmen? Ohne solche Routinen bleibt Governance theoretisch.

6.6 Dimension 6: Kompetenz, Akzeptanz und menschliche Aufsicht

KI-Systeme verändern Arbeit. Governance muss deshalb Kompetenzen und Akzeptanz berücksichtigen. Nutzerinnen und Nutzer müssen verstehen, was ein System kann, was es nicht kann, welche Daten genutzt werden dürfen, wann Ergebnisse überprüft werden müssen und wann menschliche Entscheidung erforderlich ist. Der EONAR-Beitrag zum Lernen mit KI betont, dass KI als Werkzeug verstanden werden sollte und nicht als Ersatz menschlicher Verantwortung [4].

Menschliche Aufsicht darf nicht nur formal auf dem Papier stehen. Eine Person, die Ergebnisse prüfen soll, braucht Zeit, Kompetenz, Entscheidungsmacht und Zugang zu relevanten Informationen. Wenn ein Mensch nur noch automatisch erzeugte Ergebnisse durchwinkt, entsteht keine wirksame Kontrolle. Gute Aufsicht bedeutet, dass Grenzen, Prüfpflichten und Eskalationsrechte klar sind.

Praxisregel

Ein KI-Use-Case sollte erst produktiv gehen, wenn Zweck, Daten, Risiko, Verantwortliche, Nutzerinformation, Monitoring und Stop-/Eskalationsweg dokumentiert sind.

7 Implikationen für die Umsetzung

Für die Umsetzung empfiehlt sich ein stufenweises Vorgehen. Der erste Schritt ist ein KI-Inventar. Viele Unternehmen werden feststellen, dass bereits mehr KI genutzt wird als offiziell bekannt ist: Browser-Tools, Office-Assistenten, Coding-Hilfen, Chatbots, Übersetzung, Bildgenerierung, Meeting-Zusammenfassungen oder Anbieterfunktionen in bestehenden Systemen. Dieses Inventar sollte nicht als Kontrollaktion kommuniziert werden, sondern als Grundlage für sichere Nutzung.

Der zweite Schritt ist eine einfache KI-Richtlinie. Sie sollte nicht mit Verboten beginnen, sondern mit Spielregeln: Welche Daten dürfen in externe KI-Systeme eingegeben werden? Welche Anwendungen sind frei nutzbar? Welche benötigen Freigabe? Wie müssen KI-generierte Inhalte gekennzeichnet oder geprüft werden? Wer hilft bei Fragen? Welche Mindeststandards gelten für neue Tools? Eine solche Richtlinie schafft Orientierung und reduziert Unsicherheit.

Der dritte Schritt ist ein Freigabeprozess für produktive oder risikorelevante Use Cases. Dieser Prozess sollte kurz sein, aber verbindliche Fragen stellen: Welches Problem wird gelöst? Welche Daten werden verarbeitet? Welche Entscheidungen werden beeinflusst? Welche Risiken wurden identifiziert? Welche menschliche Aufsicht ist vorgesehen? Wie wird getestet? Wie wird gemessen? Wer trägt Verantwortung?

Der vierte Schritt betrifft Pilotierung. KI-Piloten sollten nicht nur technische Machbarkeit prüfen, sondern auch Governance-Annahmen: Verstehen Nutzer die Grenzen? Sind Datenflüsse klar? Funktioniert Monitoring? Sind Ergebnisse zuverlässig genug? Werden falsche Ergebnisse erkannt? Ist der Nutzen größer als der Kontrollaufwand? Ein Pilot ohne Lernfragen liefert wenig.

Der fünfte Schritt ist die Integration in bestehende Managementsysteme. KI-Governance sollte nicht als Sonderwelt neben Datenschutz, Informationssicherheit, Qualitätsmanagement und Prozessmanagement stehen. Schnittstellen sind unvermeidlich: Datenschutz prüft personenbezogene Daten, Informationssicherheit prüft Zugriffe und Anbieter, Qualitätsmanagement prüft Prozesswirkung, Fachbereiche prüfen Nutzen und Genauigkeit, Führung priorisiert Ressourcen.

8 Diskussion

Der vorgeschlagene Rahmen versucht, eine Balance zwischen Innovationsfähigkeit und Kontrolle herzustellen. Sein Nutzen liegt darin, KI-Governance auf konkrete Use Cases herunterzubrechen. Mittelständische Organisationen benötigen selten ein großes AI-Governance-Office. Sie benötigen aber klare Mindeststandards, damit KI nicht unsichtbar, unkontrolliert oder ohne Lernschleifen eingesetzt wird.

Gegen einen solchen Rahmen wird häufig eingewandt, dass Governance KI-Projekte verlangsamt. Das kann passieren, wenn Governance als schwerfälliger Freigabeapparat gestaltet wird. Der hier vorgeschlagene Rahmen ist deshalb risikobasiert. Niedrigrisiko-Anwendungen brauchen einfache Regeln. Höhere Risiken brauchen stärkere Prüfung. Diese Differenzierung entspricht der Logik des AI Act und der Risikomanagementlogik von NIST.

Schwerer wiegt der Einwand, dass KI-Risiken technisch schwer messbar sind. NIST weist selbst auf Herausforderungen der Risikomessung hin [5]. Daraus folgt jedoch nicht, dass Organisationen auf Messung verzichten sollten. Sie können mit pragmatischen Indikatoren arbeiten: Fehlerquoten, Nutzerfeedback, Eskalationen, Datenqualitätsprobleme, manuelle Korrekturen, Zeitersparnis, Akzeptanz und Vorfälle.

Bleibt die Ressourcenfrage. KMU können nicht jede Norm und jeden Rechtsakt vollständig operationalisieren. Genau deshalb braucht es Übersetzung. ISO/IEC 42001, NIST AI RMF, OECD Principles und der AI Act liefern Orientierungslogiken. Die operative Frage lautet, welche Minimalprozesse im jeweiligen Kontext angemessen sind. Der Rahmen schlägt dafür keine maximale, sondern eine minimale Governance vor.

9 Limitationen

Dieses Working Paper ist konzeptionell und erhebt keine Primärdaten. Der vorgeschlagene Rahmen wurde nicht empirisch validiert. Er ist daher als strukturierte Praxislogik zu verstehen, nicht als Nachweis für bestimmte Erfolgswirkungen.

Der Beitrag ersetzt keine Rechtsberatung. Der EU AI Act enthält detaillierte Pflichten,

Rollen und Übergangsregelungen, deren Anwendung vom konkreten System, Einsatzkontext und Marktbezug abhängt. Organisationen sollten bei rechtlich relevanten oder potenziell hochriskanten Anwendungen fachkundige Prüfung einholen.

Die Übertragbarkeit hängt von Branche, Datenlage, Regulierung, Organisationsgröße, KI-Reife und technischer Architektur ab. Ein Unternehmen, das KI nur für interne Textentwürfe nutzt, benötigt andere Kontrollen als ein Unternehmen, das KI in Bewerbungsprozessen, Lernpfaden, Kreditentscheidungen, medizinischen Anwendungen oder sicherheitskritischen Prozessen einsetzt.

Schließlich ändern sich KI-Technologien und regulatorische Auslegung dynamisch. Governance muss deshalb als lernendes System verstanden werden. Richtlinien, Inventar, Risikokriterien und Monitoring sollten regelmäßig überprüft und an neue Anwendungen angepasst werden.

10 Fazit

KI-Governance im Mittelstand ist keine abstrakte Compliance-Übung. Sie ist die Voraussetzung dafür, KI sinnvoll, sicher und akzeptiert in Organisationen einzuführen. Wer nur Tools testet, erzeugt möglicherweise kurzfristige Effekte, aber keine belastbare Fähigkeit. Wer dagegen Zweck, Daten, Risiko, Rollen, Monitoring und Kompetenz zusammenführt, schafft eine Grundlage für produktive Nutzung.

Der Beitrag schlägt einen sechsdimensionalen Rahmen vor: Use-Case-Inventar, Risikoklassifizierung, Daten- und Modellverantwortung, Rollen und Entscheidungsrechte, Betrieb und Monitoring sowie Kompetenz, Akzeptanz und menschliche Aufsicht. Diese Dimensionen übersetzen AI Act, NIST AI RMF, ISO/IEC 42001 und OECD AI Principles in eine pragmatische Mittelstandslogik.

Der praktische Mehrwert liegt in Klarheit. Ein KI-Projekt sollte nicht mit der Frage beginnen, welches Tool am meisten verspricht, sondern mit der Frage, welcher Use Case verantwortet werden kann. KI wird dann nicht zur Black Box im Tagesgeschäft, sondern zu einer gestaltbaren Organisationsfähigkeit.

Interessenkonflikte

Der Autor erklärt, dass keine Interessenkonflikte bestehen.

Daten- und Materialverfügbarkeit

Für diesen konzeptionellen Beitrag wurden keine Primärdaten erhoben. Ergänzende Materialien sind auf der Publikationsseite verfügbar, sofern angegeben.

Quellenbasis

Dieser Beitrag basiert auf geprüften EONAR-Fachbeiträgen zu KI-Agenten und KI-gestütztem Lernen sowie auf externen Primär- und Standardisierungsquellen, insbesondere Regulation (EU) 2024/1689, NIST AI RMF 1.0, NIST Generative AI Profile, OECD AI Principles und ISO/IEC 42001:2023. Nicht auffindbare Blogtitel wurden nicht als Belege verwendet.

Autoreninformation

Lucien André Reuter ist Digital-Transformation-Experte, IT-Executive und Unternehmer mit besonderem Fokus auf digitale Geschäftsprozesse, Automatisierung, künstliche Intelligenz und technologiegestützte Organisationsentwicklung. Als Geschäftsführer der EONAR GmbH begleitet er Unternehmen bei der strategischen und operativen Weiterentwicklung digitaler Plattformen, Prozesslandschaften und KI-gestützter Automatisierungslösungen.

Zuvor verantwortete er als Chief Technology Officer und Head of Digital Enablement an der DTMD University for Digital Technologies in Medicine and Dentistry in Luxemburg den Aufbau und Betrieb einer integrierten IT-, Lern- und Prozesslandschaft. Seine Arbeit verbindet technologische Architektur, Governance, Nutzerakzeptanz und Effizienzorientierung in komplexen Organisationen.

Neben seiner unternehmerischen Tätigkeit engagiert sich Lucien André Reuter in Lehre, Wissenstransfer und öffentlicher Fachkommunikation. Als Host des Podcasts „Digital Sense – The Human Side of AI“ spricht er mit Gästen über die menschliche, organisatorische und strategische Dimension von künstlicher Intelligenz. Im Mittelpunkt stehen dabei nicht nur technologische Möglichkeiten, sondern vor allem die Frage, wie Menschen, Führungskräfte und Organisationen KI sinnvoll, verantwortungsvoll und wirksam einsetzen können.

Akademisch verbindet Lucien André Reuter Informatik, Management und angewandte Digitalisierung. Er verfügt über einen Bachelor of Science in Informatik, einen MBA sowie einen Doctor of Business Administration mit Forschungsschwerpunkt Digitalisierung und Prozessoptimierung in Organisationen. Sein Profil verbindet Praxis, Forschung, Lehre und mediale Fachkommunikation an der Schnittstelle von digitaler Transformation, KI und organisationalem Wandel.

References

- [1] Eisenhardt, K. M. (1989). Building theories from case study research. *The Academy of Management Review*, 14(4), 532–550. <https://doi.org/10.2307/258557>
- [2] European Parliament and Council of the European Union. (2024). *Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence*. Official Journal of the European Union. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- [3] ISO/IEC. (2023). *ISO/IEC 42001:2023: Information technology – Artificial intelligence – Management system*. International Organization for Standardization. <https://www.iso.org/standard/81230.html>
- [4] Müller, K. M., & Grzeszkowiak, J.-L. (2025). *Lernen mit Künstlicher Intelligenz: Chancen, Herausforderungen und Gestaltungsperspektiven*. EONAR GmbH, 23. Juni 2025. <https://www.eonar.de/lernen-mit-kuenstlicher-intelligenz/>
- [5] National Institute of Standards and Technology. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. NIST AI 100-1. <https://doi.org/10.6028/NIST.AI.100-1>
- [6] National Institute of Standards and Technology. (2024). *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*. NIST AI 600-1. <https://doi.org/10.6028/NIST.AI.600-1>
- [7] OECD. (2024). *OECD AI Principles overview*. OECD.AI Policy Observatory. <https://oecd.ai/en/ai-principles>
- [8] Reuter, L. A. (2025). *Wie man einen AI Agent baut – Schritt für Schritt erklärt*. EONAR GmbH, 5. November 2025. <https://www.eonar.de/wie-man-einen-ai-agent-baut-schritt-fuer-schritt-erklaert/>