



Dualer
CONSULTING

| AI Agents Whitepaper

Autoren und Projektpartner



Philipp Weissgerber

Dualer Consulting e.V.
Projektleiter
Mitwirkung an Use Case 1



Felix Nitsch

Dualer Consulting e.V.
Projektteammitglied
Mitwirkung an Use Case 1



David Bruckmeier

Dualer Consulting e.V.
Projektteammitglied
Mitwirkung an Use Case 2



Cedric Korosec

Dualer Consulting e.V.
Projektteammitglied
Mitwirkung an Use Case 2



Nils Michelsen

Dualer Consulting e.V.
Projektteammitglied
Mitwirkung an Use Case 3



Mathavan Ganeshanathan

Dualer Consulting e.V.
Projektteammitglied
Mitwirkung an Use Case 3



1. Management Summary

Dieses Whitepaper zeigt anhand von drei konkreten Anwendungsfeldern, jeweils von unterschiedlichen Projektteams erarbeitet, wie agentische AI in komplexen, wissensintensiven Unternehmensprozessen eingesetzt werden kann und welchen strategischen Mehrwert sie erzeugt. Die unterschiedlichen Perspektiven ermöglichen dabei eine bewusst vielfältige Betrachtung agentischer Systemarchitekturen.

Im ersten Use Case wird ein IoT-gestütztes Abweichungsmanagement in asset-intensiven Industrieumgebungen betrachtet. Ziel ist es, aus einer Vielzahl technischer Warnmeldungen priorisierte, managementtaugliche Fallkarten zu erzeugen. Durch die Kombination aus kontextueller Bewertung, Wirkungsabschätzung und strukturiertem Wissensmanagement werden Reaktionszeiten verkürzt, Entscheidungsqualität erhöht und Erfahrungswissen systematisch skalierbar gemacht.

Der zweite Use Case fokussiert auf einen Management Decision Agent, der fragmentierte Entscheidungsgrundlagen konsolidiert, KPIs harmonisiert und priorisierte Entscheidungsoptionen bereitstellt. Die agentische Architektur ermöglicht eine proaktive, transparente und nachvollziehbare Managementsteuerung, reduziert kognitive Belastung und erhöht die Konsistenz strategischer Entscheidungen, ohne die menschliche Entscheidungshoheit zu ersetzen.

Im dritten Use Case wird ein agentisches Debitoren- und Cash-Management beschrieben. Durch die Integration von Transaktionsdaten, Bonitätsinformationen und kontextuellen Kommunikationsdaten entsteht ein System, das Zahlungsrisiken frühzeitig erkennt, Mahnprozesse intelligent steuert und Liquidität in Echtzeit absichert. Der Fokus liegt dabei nicht nur auf Effizienz, sondern auf einer balancierten Steuerung von Risiko, Kundenbeziehung und Working Capital.

Über alle drei Use Cases hinweg wird deutlich: Der eigentliche Mehrwert agentischer AI liegt nicht in isolierter Automatisierung, sondern in der strukturierten End-to-End-Architektur, die Daten, Kontext, Wirkung und Handlungsempfehlungen systematisch verbindet. Agentische Systeme erhöhen Transparenz, Priorisierungsfähigkeit und Reproduzierbarkeit von Entscheidungen, bei klar definierter Governance und verpflichtender menschlicher Freigabe in kritischen Situationen.

Dieses Whitepaper vermittelt ein grundlegendes Verständnis von agentischer AI und Multi-Agenten-Systemen im Beratungskontext. Anhand der drei ausgewählten Use Cases wird gezeigt, wie diese Systemlogik praktisch angewendet werden kann. Die Use Cases dienen als konzeptioneller Umsetzungsansatz und werden aus der Perspektive der studentischen Unternehmensberatung Dualer betrachtet.



AI Agents Whitepaper

1. Management Summary	2
2. Einleitung.....	4
3. Use Case 1 IOT Abweichungsmanagement	5
3.1 Pain Points	5
3.2 Technische Umsetzung	8
3.3 Rechtliche Umsetzung.....	21
3.4 Abschließendes Fazit zum Use Case IOT Abweichungsmanagement	24
4. Use Case 2 Management Design Decision Agent	25
4.1 Pain Points	25
4.2 Technische Umsetzung	27
4.3 Rechtliche Umsetzung.....	31
4.4 Abschließendes Fazit zum Use Case Management Decision Agent	33
5. Use Case 3 Debitorenmanagement	34
5.1 Pain Points	34
5.2 Technische Umsetzung	36
5.3 Rechtliche Umsetzung.....	47
5.4 Abschließendes Fazit zum Use Case Debitorenmanagement	48
6. Handlungsempfehlung	49
Definitionsverzeichnis.....	50



2. Einleitung

Künstliche Intelligenz wird in Unternehmen bereits vielfach zur Effizienzsteigerung und Entscheidungsunterstützung eingesetzt. Der Fokus liegt dabei häufig auf klar abgegrenzten Anwendungsfällen wie Datenanalysen, Prognosen oder der Automatisierung einzelner Prozessschritte. Diese Systeme agieren überwiegend reaktiv innerhalb vordefinierter Abläufe und steuern Prozesse nicht ganzheitlich oder zielorientiert. Agentische AI erweitert diesen Ansatz grundlegend. Darunter werden AI-Systeme verstanden, die eigenständig Informationen aus unterschiedlichen Quellen erfassen, diese einordnen und auf Grundlage klar definierter Zielsetzungen Handlungsvorschläge ableiten oder Aufgaben selbstständig ausführen können. Im Unterschied zu klassischer AI verfolgen agentische Systeme Aufgaben proaktiv, priorisieren Informationen und passen ihr Vorgehen an veränderte Rahmenbedingungen an. Sie agieren zielorientiert, kontextbezogen und kontinuierlich, stets innerhalb klar definierter Leitplanken.

In der Praxis handelt es sich dabei meist nicht um einen einzelnen AI-Agenten, sondern um sogenannte Multi-Agenten-Systeme. Diese bestehen aus mehreren spezialisierten Agenten, die jeweils klar abgegrenzte Aufgaben übernehmen und ihre Ergebnisse untereinander austauschen. Durch diese koordinierte Zusammenarbeit lassen sich komplexe, wissensintensive Prozesse strukturierter und effizienter unterstützen.

Für Unternehmensberatungen eröffnet agentische AI neue Einsatzmöglichkeiten, sowohl zur internen Effizienzsteigerung als auch als Bestandteil innovativer Beratungsleistungen. Multi-Agenten-Systeme unterstützen Beraterinnen und Berater bei der Analyse großer Informationsmengen, der strukturierten Aufbereitung von Entscheidungsgrundlagen und der Begleitung fundierter Entscheidungsprozesse. Die Verantwortung für Entscheidungen verbleibt dabei bewusst beim Menschen; agentische AI unterstützt, strukturiert und analysiert.

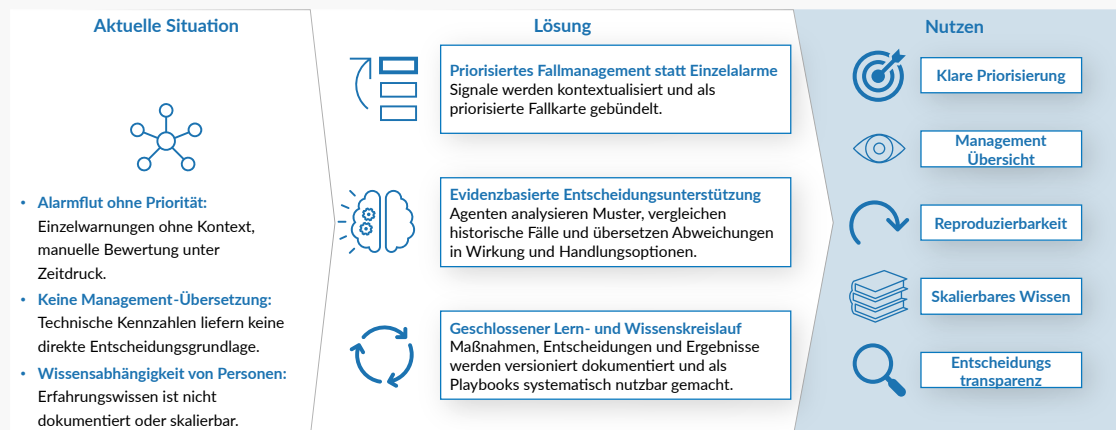


3. Use Case 1 | IoT Abweichungsmanagement

In asset-intensiven Industrieumgebungen sollen Abweichungen aus Sensorik und Leitsystemen so verarbeitet werden, dass aus vielen Einzelmeldungen wenige, priorisierte Bearbeitungsfälle entstehen. Im Fokus stehen Produktionslinien, Maschinen oder Netzsegmente, in denen Warnmeldungen heute operativ binden, aber selten direkt entscheidungsfähig sind. Die folgenden Pain Points beschreiben die operativen Engpässe, die im anschließenden Kapitel „Technische Umsetzung“ systematisch adressiert werden.

Operative Entscheidungen werden unterstützt, während eine skalierbare Wissensbasis entsteht

Use Case: IOT Abweichungsmanagement



3.1 Pain Points

Pain Point 1: Zu viele Warnmeldungen ohne klare Priorität

Im IoT-gestützten Abweichungsmanagement von Industrie-, Energie- und Versorgungsunternehmen entstehen täglich eine Vielzahl von Warn- und Abweichungsmeldungen. Anlagen, Netze und Sensoren überwachen kontinuierlich technische Zustände und melden Grenzwertüberschreitungen, Auffälligkeiten oder Abweichungen vom Sollbetrieb.

Das heutige Abweichungsmanagement ist dabei überwiegend auf die Erfassung einzelner technischer Signale ausgelegt. Warnmeldungen werden typischerweise isoliert in Dashboards, Listen oder Alarmhistorien dargestellt. Eine systematische Einordnung nach Relevanz, Dringlichkeit oder potenzieller Auswirkung findet häufig nicht statt. In der Praxis bedeutet dies, dass eine Vielzahl an Warnmeldungen gleichwertig nebeneinander steht, unabhängig davon, ob es sich um harmlose Schwankungen oder um Vorboten kritischer Störungen handelt.

Der entscheidende Engpass liegt in der fehlenden kontextuellen Bewertung. Mitarbeitende müssen Warnmeldungen manuell sichten, vergleichen und priorisieren. Dabei fehlt häufig der Gesamtzusammenhang, etwa der zeitliche Verlauf, die Wechselwirkung mehrerer Sensoren oder der Vergleich mit ähnlichen historischen Situationen. Kritische Entwicklungen werden dadurch entweder übersehen oder erst dann erkannt, wenn sich das Problem bereits deutlich verschärft hat. Entscheidungen erfolgen reaktiv und unter Zeitdruck.



Für die Praxis ergibt sich folgendes Bild: Operative Teams sind stark mit der Prüfung und Bewertung von Warnmeldungen beschäftigt, ohne dass klar erkennbar ist, welche Abweichung tatsächlich unmittelbaren Handlungsbedarf erfordert. Die Fähigkeit zur Priorisierung hängt stark von individueller Erfahrung und verfügbarer Kapazität ab. In Phasen hoher Alarmdichte steigt das Risiko, relevante Abweichungen zu spät oder gar nicht zu adressieren.

Genau an dieser Stelle kann ein AI-Agent im IoT-gestützten Abweichungsmanagement ansetzen. Er übernimmt die kontextuelle Einordnung und Priorisierung von Warnmeldungen. Durch den Abgleich aktueller Abweichungen mit historischen Mustern, ähnlichen Fällen und bekannten Fehlerbildern ist ein AI-Agent in der Lage, Warnmeldungen automatisch nach Relevanz, Dringlichkeit und potenzieller Auswirkung zu bewerten. Statt einer Vielzahl gleichwertiger Alarme erhalten Mitarbeitende eine fokussierte Sicht auf jene Abweichungen, die tatsächlich entscheidungsrelevant sind. Dadurch werden Reaktionszeiten verkürzt, kritische Entwicklungen früher erkannt und operative Entscheidungen systematisch unterstützt.

Pain Point 2: Mangelnde Managementtauglichkeit

Im produzierenden Gewerbe fokussieren sich täglich entstehende Abweichungen zu großen Teilen auf technischen Kennzahlen. Diese stellen keine unmittelbare, managementtaugliche Entscheidungsgrundlagen dar.

Das moderne Abweichungsmanagement auf Basis von IoT- und Produktionsdaten konzentriert sich heute primär auf operative, technische Größen wie: Laufzeiten, Stillstände, Ausschussquoten, Energieverbräuche oder Grenzwertverletzungen. Diese technischen Informationen sind für Instandhaltung und Produktion wertvoll, für das Management jedoch nur eingeschränkt nutzbar.

Der entscheidende Engpass liegt in der fehlenden Transferleistung. Bevor technische Abweichungen als managementtaugliche Entscheidungsgrundlage genutzt werden können müssen diese manuell in wirtschaftliche Aussagen übersetzt werden. Finanzielle Effekte, etwa entgangene Produktionsmengen bei Stillständen, zusätzliche Personalkosten, Vertragsstrafen durch verspätete Lieferungen oder langfristige Qualitätseinbußen, müssen manuell dem Management aufbereitet werden. Selbes gilt für strategische Folgen wie Risiken für Kundenbeziehungen, Liefertreue oder Kapazitätsplanung. Diese manuelle Übersetzung ist zeitaufwendig, fehleranfällig und zudem stark von individuellem Know-how abhängig. Für die Praxis ergibt sich folgendes Bild: Das Management erhält zwar Berichte über technische Abweichungen, kann diese jedoch nicht unmittelbar für Entscheidungen nutzen. Konkrete finanzielle und strategische Auswirkungen können erst nach manueller Überarbeitung in Meetings, Excel-Analysen oder durch Erfahrungswissen einzelner Experten, dem Management zur Verfügung gestellt werden. In der Folge sind Handlungsmaßnahmen nicht nur stark reaktiv, sondern verstärkt zeitlich verzögert.

Genau an diesen Transfer könnte ein AI-Agent im IoT-gestützten Abweichungsmanagement ansetzen. Er übernimmt die zentrale Transferleistung von technischer Realität hin zu betriebswirtschaftlicher Steuerung. Neben IoT-Daten und Produktionskennzahlen werden Kostenmodelle und Planungsdaten durch einen AI-Agent analysiert. Die Ausgabe stellt der automatisierte Transfer einer technischen Abweichung in konkrete monetäre und strategische Auswirkungen dar. Möglich wären folgende Vorstellungen: Stillstände werden direkt als entgangener Deckungsbeitrag quantifiziert, Qualitätsabweichungen als potenzielle Reklamations- oder Nacharbeitskosten bewertet. Darüber hinaus kann ein AI-Agent qualitative und strategische Folgen unterschiedlicher Abweichungen, etwa Auswirkungen auf



Liefertermintreue, Kundenprioritäten oder Engpassressourcen analysieren und bewerten. Aus diesem Transfer lassen sich unmittelbare und reaktionsschnelle Entscheidungen für das Management ableiten. So kann es klar priorisieren, welche Maschine, Linie oder Abweichung sofortige Aufmerksamkeit erfordert, weil sie den größten finanziellen oder strategischen Schaden verursacht.

Pain Point 3: Fehlendes Wissensmanagement

In etablierten Unternehmen hängen Entscheidungen im Abweichungsmanagement in großer Zahl von einzelnen, erfahrenen Mitarbeitenden ab. Systematische, skalierbare Entscheidungsgrundlagen existieren kaum. Erfahrungswissen ist nicht dokumentiert, Abweichungen sind nicht vergleichbar und Reaktionen in der Folge inkonsistent.

Das heutige, gängige Abweichungsmanagement basiert regelmäßig auf ad-hoc Entscheidungen und individueller Expertise. Erfahrungswissen liegt bei wenigen, oft sehr erfahrenen Mitarbeitenden, wird nicht systematisch dokumentiert und ist nicht für Nachfolger, andere Schichten oder Werke verfügbar.

Dies führt zu inkonsistenten Reaktionen. Dieselbe Abweichung wird je nach Person oder Abteilung unterschiedlich bewertet, priorisiert und behandelt. Darüber hinaus ist eine Standardisierung in der Dokumentation und Nachverfolgung von Maßnahmen keine gängige Praxis. Abweichungen werden abteilungsspezifisch und werksübergreifend, personenindividuell und unterschiedlich deklariert. Maßnahmen werden nicht konsequent erfasst, wodurch ihre Wirksamkeit häufig im Unklaren bleibt. In der Folge lassen sich Best Practices nicht skalieren und Abweichungen nicht vergleichen. Beschriebenes lässt systematisches Lernen aus Abweichungen nicht stattfinden. Für das Management bedeutet dies: Es fehlt Transparenz darüber, welche Maßnahmen ergriffen wurden, ob sie erfolgreich waren und welche Abweichungen einen ähnlichen Handlungsbedarf erfordern würden. Interne Benchmarks oder Kennzahlen, um Abweichungen werksübergreifend vergleichbar zu machen, um daraus Zielsetzungen ableiten und steuern zu können, sind nicht existent.

Das Wissensmanagement könnte ein AI-Agent im IoT-gestützten Abweichungsmanagement übernehmen. Er macht vorhandenes Erfahrungswissen reproduzierbar und stellt es allen Mitarbeitenden konsistent zur Verfügung. Gleiche oder sich ähnelnde Abweichungen werden automatisch identifiziert, mit Best Practices verknüpft und auf durchgeführte Maßnahmen verwiesen. Diese Möglichkeit besteht über Abteilungs- und Werksgrenzen hinweg. Zudem ermöglicht der AI-Agent einen systematischen Vergleich von Abweichungen, woraus sich interne Benchmarks ergeben. Dabei identifiziert der AI-Agent Wiederholungsmuster clustert diese und ist in der Lage hieraus standardisierte Handlungsempfehlungen abzuleiten. Gleichzeitig werden Maßnahmen automatisch erfasst, klar dokumentiert und ihr Status nachverfolgt, sodass Lernprozesse aus Abweichungen kontinuierlich stattfinden. Durch diesen Ansatz hängen Entscheidungen weniger von Subjektivität und persönlicher Erfahrung ab. Das Management erhält klare Transparenz über Maßnahmenumsetzung und Wirksamkeit und kann daraus schnelle, fundierte, unabhängige Entscheidungen treffen.



3.2 Technische Umsetzung

Zielsetzung

Dieses Kapitel beschreibt die technische Umsetzung eines IoT-gestützten Abweichungsmanagements in asset-intensiven Industrieumgebungen wie Fertigung, Energieerzeugung und Versorgungsnetzen. Unter IoT wird hier die vernetzte Erfassung technischer Zustände verstanden. Diverse Sensoren liefern kontinuierlich Messwerte, Leit- sowie Betriebssysteme erzeugen Ereignisse wie Alarme, Zustandswechsel und Qualitätsmeldungen.

Das betrachtete Setting ist bewusst konkret beschrieben: Im produzierenden Umfeld sind typische Betrachtungseinheiten eine Produktionslinie, eine Maschine oder ein Aggregat (zum Beispiel Pumpen, Kompressoren, Motoren, Getriebe, Fördertechnik). In Energie und Versorgungsnetzen sind es Netzsegmente oder Netzelemente (zum Beispiel Trafostationen, Leitungsabschnitte, Schaltanlagen, Erzeugungsaggregate). Typische Sensorarten sind unter anderem Temperatur, Druck, Vibration, Stromaufnahme, Durchfluss, Drehzahl, Ventilstellung und Spannung.

Typische Meldungsklassen, die in der Praxis hohe Volumina erzeugen, sind:

- Grenzwertmeldungen (Wert über/unter Grenzbereich)
- Trend- und Driftmeldungen (schleichende Verschlechterung über Zeit)
- Zustandswechsel (Start/Stop, Betriebsmoduswechsel, Lastwechsel)
- Qualitätsmeldungen (Ausschussindikatoren, Messmittelwarnungen, Prozessfähigkeitsverletzungen)
- Kommunikations- und Datenqualitätsmeldungen (Signalverlust, Zeitstempelabweichung, Sensorfehler)

Die Architektur wird entlang eines durchgehenden Nutzungsfalls erklärt, der sich an der Realität auf dem Shopfloor orientiert: Leitstand und Schichtführung benötigen eine konzentrierte Sicht auf wenige wirklich relevante Abweichungen. Hingegen braucht die Instandhaltung ein belastbares technisches Gesamtbild, inklusive Vergleich, mit ähnlichen historischen Situationen. Die Produktionsleitung benötigt eine managementtaugliche Einordnung in Wirkung, Dringlichkeit und Handlungsoptionen.

Jede technische Ebene wird so begründet, dass ihr Beitrag zu den drei Engpässen nachvollziehbar ist. Erstens muss die hohe Alarm- und Warnmeldungsichte in eine priorisierte Bearbeitungslogik überführt werden. Zweitens müssen technische Abweichungen in eine managementtaugliche Entscheidungsgrundlage übersetzt werden. Drittens muss Erfahrungswissen aus Einzelköpfen in skalierbares, standortübergreifendes Wissen überführt werden.

Abgrenzung: Entscheidungsunterstützung statt automatisiertem Eingriff

Das System ist als Entscheidungsunterstützung konzipiert. Es erzeugt priorisierte, nachvollziehbar begründete Fallkarten und Handlungsvorschläge. Automatisierte Eingriffe in den Anlagenbetrieb sind nicht Gegenstand dieser Architektur. Entscheidungen mit potenzieller Sicherheits- oder Produktionswirkung (zum Beispiel Stop, Bypass, De-rate, Umschaltung) erfordern verpflichtend eine menschliche Freigabe.



Erwarteter Output im Alltag

Der erwartete, sichtbare Output ist keine zusätzliche Alarm-Liste, sondern eine Fallkarte (im Ticketing, im Leitstand oder im Instandhaltungssystem), die konsistent die folgenden Elemente enthält:

- Priorität (begründet, nicht nur Alarmstufe)
- Kontext (Betriebsmodus, Schicht, Auftrag/Netzzustand, Ereigniskette)
- Wirkung (managementtaugliche Einordnung als Bandbreite, mit Annahmen)
- Maßnahmen (abgestuft und prüfbar)
- Begründung (Evidenzverweise, Unsicherheit, offene Punkte)

Ausgangssituation & Leitfall

Im Tagesgeschäft beginnt das Problem häufig als eine lange Liste von Warnmeldungen im Leitstand, ergänzt um Meldungen aus Instandhaltungssystemen. Ein einzelner Alarm beantwortet selten die entscheidende Frage: Ist das nur ein kurzzeitiges Rauschen oder der Vorbote einer Störung, die Produktion, Qualität oder Liefertermine gefährdet.

Die technische Umsetzung folgt deshalb einem zentralen Prinzip: Ungewöhnlich ist nicht automatisch bearbeitungsrelevant. Der technische Kern ist eine Kette, die aus vielen Rohsignalen wenige Bearbeitungsfälle erzeugt. Ein Bearbeitungsfall ist eine strukturierte Fallkarte, die für Technik und Management gleichermaßen verständlich ist, weil sie das technische Muster, den betrieblichen Kontext, eine begründete Priorität, eine Wirkungseinordnung und konkrete nächste Schritte zusammenführt.

Um die Logik durchgängig greifbar zu machen, wird im Folgenden ein Leitfall verwendet:

In einer Produktionslinie zeigt ein Antriebsmotor mit Lager über mehrere Stunden eine schleichend steigende Vibration, gleichzeitig steigt die Stromaufnahme leicht an. Kurz darauf treten kurze Mikro-Stopps auf, und es erscheint eine Qualitätswarnung (Ausschussindikator). Im Leitstand entstehen dazu mehrere Warnmeldungen aus unterschiedlichen Quellen, die isoliert betrachtet harmlos wirken könnten.

Dieses Leitbeispiel ist typisch: Relevanz entsteht nicht durch einen Einzelalarm, sondern durch die Kombination aus Verlauf, Mehrsensorik, Betriebsmodus und Historie.

3.2.1 Architektur

Die Architektur ist in Ebenen gegliedert, weil sich Datenaufnahme, Datenaufbereitung, Erkennung, Fallsteuerung und Entscheidungsunterstützung in ihrer Logik unterscheiden. Ohne diese Trennung entstehen in der Praxis zwei typische Probleme: Die Alarmflut wird technisch verstärkt, weil Rohsignale direkt in Tickets übersetzt werden, und die Begründbarkeit leidet, weil später nicht nachvollziehbar ist, welche Verarbeitungsschritte zu einem Fall geführt haben.

Im Folgenden wird jede Ebene als Baustein einer End-to-End-Geschichte erklärt. Nach jedem Baustein steht ein Wirkungssatz, der die Verbindung zu den drei Engpässen explizit herstellt.



Ebene 1: Sensor und Edge, Stabilität vor Intelligenz

In industriellen Anlagen und Netzen sind Sensordaten nie vollständig ideal. Messwerte können rauschen (zufällige Schwankungen), Sensoren können driften (langsame Verschiebung des Messniveaus), und Kommunikation kann ausfallen. Wenn diese Realität nicht früh berücksichtigt wird, entstehen Fehlalarme oder scheinbare Abweichungen, die rein datengetrieben sind und das operative Team zusätzlich belasten.

Die Edge-Ebene ist deshalb kein Ort für komplexe Entscheidungen, sondern für Robustheit. Technisch bedeutet das, dass Messwerte vor der Weiterleitung plausibilisiert, zeitlich sauber gestempelt und bei Verbindungsproblemen gepuffert werden. Ebenso werden Datenqualitätskennzeichen erzeugt, die später sichtbar machen, ob ein Signal vertrauenswürdig ist. Ein Datenqualitätskennzeichen ist eine Markierung, die zum Beispiel anzeigt, ob Werte fehlen, ob es Ausreißer gab oder ob die Zeitbasis unsicher ist.

Im Leitbeispiel bedeutet das: Vibration und Stromaufnahme werden bereits an der Quelle so aufbereitet, dass Messaussetzer oder Zeitversatz nicht als „scheinbare“ Abweichung im Leitstand erscheinen.

Diese Ebene reduziert die Zahl falscher oder irreführender Warnmeldungen, weil Datenqualität früh sichtbar gemacht wird. Sie verbessert die managementtaugliche Einordnung indirekt, weil spätere Übersetzungen auf verlässlichen Messgrundlagen basieren. Sie schafft zudem die Voraussetzung für standortübergreifendes Wissensmanagement, weil nur saubere und vergleichbare Messungen sinnvoll in historischen Fallpools landen.

Ebene 2: Data Ingestion, kontrollierte Aufnahme statt Systemwildwuchs

Ein IoT-Abweichungsmanagement verknüpft typischerweise verschiedene Quellen: Sensorströme, Alarmereignisse aus Leitsystemen, Betriebsdaten aus Produktionssystemen sowie Stammdaten und Instandhaltungsdaten. Zusätzlich integriert der Ingestion-Layer externe, nicht-sensorische Quellen wie Confluence- und SharePoint-Dokumente, Wartungsprotokolle, PDF-Berichte oder von Mitarbeitenden erstellte Inhalte. Diese externen Daten werden als separate Wissensobjekte erfasst, validiert und eindeutig den betroffenen Assets zugeordnet. Wenn diese Quellen Punkt-zu-Punkt integriert werden, entsteht eine fragile Landschaft, bei der jede Quellsystemänderung Nebenwirkungen im Gesamtsystem auslöst.

Der Ingestion-Layer ist deshalb ein kontrollierter Eingang, der Datenformate prüft, Quellen eindeutig zuordnet und den Datenfluss protokolliert. Unter Datenvertrag wird hier eine explizite Vereinbarung verstanden, welche Felder ein Datensatz enthalten muss und welche Bedeutung sie haben. Unter eindeutiger Asset-Zuordnung wird verstanden, dass ein Sensorwert oder Ereignis eindeutig einer Anlage oder einem Netzelement zugeordnet werden kann, auch wenn verschiedene Systeme unterschiedliche Kennungen verwenden.

Im Leitbeispiel bedeutet das: Vibration, Stromaufnahme, Mikro-Stopps und Qualitätswarnung werden als zusammengehörige Signale derselben Maschine/Station erkennbar, statt als vier getrennte Meldungen ohne gemeinsamen Bezug.

Diese Ebene reduziert Warnmeldungsdichte, die durch fehlerhafte oder inkonsistente Daten entsteht, weil sie Validierung und Konsistenz erzwingt. Sie unterstützt die managementtaugliche



Entscheidungsgrundlage, weil spätere Übersetzungen nur dann glaubwürdig sind, wenn Herkunft und Zeitpunkt der Daten nachvollziehbar sind. Sie ermöglicht skalierbares Wissensmanagement, weil nur konsistente Asset-Identitäten standortübergreifende Vergleiche zulassen.

Ebene 3: Data Processing und Feature Engineering, das Gesamtbild entsteht

Der Kern der Alarmflut liegt häufig nicht in zu vielen echten Störungen, sondern in zu vielen isolierten Einzelbeobachtungen. Ein Mensch sieht dann einzelne Grenzwertverletzungen, aber nicht, ob ein Trend kippt, ob mehrere Sensoren gemeinsam auffällig sind oder ob das Verhalten im aktuellen Betriebsmodus sogar normal ist.

Die Verarbeitungsebene bereitet Rohdaten so auf, dass Muster statt Einzelpunkte sichtbar werden. Ein Muster kann zum Beispiel ein langsamer Temperaturanstieg, ein wiederkehrender Vibrationsimpuls oder eine Kombination aus Druckabfall und Stromanstieg sein. Feature Engineering bedeutet hier, dass aus Rohmesswerten Merkmale berechnet werden, die Muster quantitativ beschreiben, etwa Trendstärke, Varianz, Sprungstellen oder gleitende Mittelwerte über definierte Zeitfenster.

Entscheidend ist außerdem die Kontextanreicherung. Kontext bedeutet betriebliche Rahmenbedingungen, etwa Auftrag, Produktvariante, Schicht, Betriebsmodus oder geplante Auslastung. Ohne diesen Kontext ist eine Abweichung häufig nicht interpretierbar. Ein hoher Energieverbrauch kann im Hochlastbetrieb normal sein und im Leerlauf kritisch.

Im Leitbeispiel bedeutet das: Das System erkennt nicht nur „Vibration hoch“, sondern „Vibration steigt kontinuierlich im stabilen Betriebsmodus“, gemeinsam mit leicht steigender Stromaufnahme und auftretenden Mikro-Stopps.

Diese Ebene löst die fehlende Priorisierung, weil sie aus Einzelalarmen ein verständliches technisches Gesamtbild erzeugt. Sie bildet die Grundlage für managementtaugliche Übersetzung, weil nur korrekt interpretierte technische Abweichungen sinnvoll in Wirkung übertragen werden können. Sie unterstützt Wissensmanagement, weil kontextualisierte Merkmale spätere Ähnlichkeitssuchen und Clustering erst belastbar machen.

Ebene 4: Abweichungserkennung

Abweichungserkennung ist die Fähigkeit, ungewöhnliche Muster zu identifizieren. Ungewöhnlich bedeutet dabei nicht zwingend falsch oder kritisch, sondern statistisch oder regelbasiert bemerkenswert im Vergleich zu einem erwarteten Verhalten. Erwartetes Verhalten kann ein Grenzwert, eine Baseline aus historischer Normalfahrt oder ein Modell sein, das typische Muster pro Betriebsmodus kennt.

Ein zentraler Designentscheid ist, dass die Erkennung keine direkte operative Arbeit auslöst, sondern zunächst Kandidaten erzeugt. Ein Kandidat ist ein Vorschlag, dass ein Muster auffällig ist, inklusive einer Bewertung und einem Evidenzpaket. Ein Evidenzpaket ist die Zusammenstellung der relevanten Sensorsegmente, Ereignisse und Kontextinformationen, die zeigen, warum der Kandidat entstanden ist.

Die Abweichungserkennung erzeugt keine Diagnose, sondern eine formal begründete Auffälligkeits-Hypothese, deren Evidenz explizit im Kandidatenobjekt gespeichert wird.



Im Leitbeispiel entsteht aus mehreren Warnmeldungen ein Kandidat wie „mögliche schleichende Lagerdegradation“. Dazu werden die relevanten Zeitfenster der Vibration, der begleitenden Stromaufnahme sowie der erkannten Mikro-Stopps zu einer konsistenten Ereigniskette zusammengeführt.

Diese Ebene reduziert die ungeordnete Warnmeldungsichte, weil sie Rohalarme in bewertete/priorisierte Kandidaten überführt und damit Rauschen ausfiltert, ohne sofort Tickets zu erzeugen. Sie unterstützt managementtaugliche Entscheidungsgrundlagen indirekt, weil Kandidaten bereits eine nachvollziehbare Evidenzstruktur enthalten. Sie schafft die Vorstufe für skalierbares Wissensmanagement, weil Kandidaten standardisiert abgelegt und später zu Fällen gebündelt werden können.

Ebene 5: Case Management, aus Signalen werden steuerbare Fälle

Operative Teams managen im Alltag keine Sensorpunkte, sondern Vorgänge. Ein Vorgang muss eine eindeutige Identität, einen Status, eine verantwortliche Rolle und eine Frist haben. Genau das leistet das Case Management.

Technisch bündelt Case Management Kandidaten, die zeitlich und inhaltlich zusammengehören. Deduplikation bedeutet, dass wiederholte Kandidaten desselben Musters nicht als neue Arbeit erscheinen, sondern im gleichen Fall zusammenlaufen. Clustering bedeutet, dass ähnliche Fälle gruppiert werden, um Wiederholungen und standortübergreifende Muster sichtbar zu machen. Ein Statusmodell beschreibt den Lebenszyklus eines Falls, zum Beispiel offen, in Prüfung, mitigiert und gelöst.

Im Leitbeispiel bedeutet das: Mehrere Vibration- und Stromaufnahme-Warnungen sowie Mikro-Stopps werden nicht als gesonderte, individuelle Tickets geführt, sondern als eine Fallkarte „Antrieb Station X zeigt progressives Auffälligkeitsmuster“.

Diese Ebene löst die hohe Alarmdichte im operativen Alltag, weil sie aus vielen Einzelmeldungen wenige Fallkarten macht, die man bearbeiten kann. Sie bildet eine Grundlage für managementtaugliche Entscheidungsgrundlagen, weil ein Fall die geeignete Einheit ist, um Wirkung, Dringlichkeit und Optionen zusammenzuführen. Sie adressiert fehlendes Wissensmanagement direkt, weil Fälle die wiederverwendbare Einheit sind, in der Maßnahmen und Ergebnisse dauerhaft dokumentiert werden.

Ebene 6: AI-Agentenebene, Entscheidungsunterstützung als kontrollierter Prozess

An dieser Stelle entsteht der wesentliche Mehrwert über klassische IoT-Analytik hinaus. Es geht nicht darum, Entscheidungen zu automatisieren, sondern Entscheidungen reproduzierbar vorzubereiten. Reproduzierbar bedeutet, dass zwei Personen bei gleicher Datenlage zu ähnlich begründeten Ergebnissen gelangen, unabhängig von individueller Erfahrung oder Tagesform.

Input:

Um unbelegte Schlussfolgerungen zu verhindern, arbeiten Agenten ausschließlich auf freigegebenen Datenobjekten:

- Fallkarte (Deviation Case) als zentrales Arbeitsobjekt
- Sensorsegmente (Zeitfenster der relevanten Signale)



- Ereignisketten (zeitlich geordnete Ereignisse wie Stopps, Zustandswechsel, Qualitätsmeldungen)
- Kostenmodell-Snapshot (versionierte Annahmen für Wirkungsabschätzung)
- Playbooks (dokumentierte Vorgehensweisen und historische Maßnahmen)

Output:

Agentenergebnisse werden nicht als freier Text abgelegt, sondern strukturiert als Bestandteil der Fallkarte:

- Priorität (inklusive Begründung)
- Wirkung (Bandbreite, Treiber, Annahmen)
- Maßnahmen (abgestuft, mit Voraussetzungen und Risiken)
- Begründung (Evidenzverweise, Unsicherheit, offene Punkte)

Nachvollziehbarkeit und Konfliktmanagement in der Agentenebene

Wenn sich Ergebnisse widersprechen (zum Beispiel: hoher technischer Schweregrad, aber geringe wirtschaftliche Wirkung; oder ähnliche Fälle widersprechen in der Maßnahmeneffektivität), markiert der Orchestrator den Konflikt explizit in der Fallkarte, fordert gezielt zusätzliche Evidenz an (zum Beispiel weiteres Sensorfenster, andere Betriebsmodi, zusätzliche Ereignisdaten) und eskaliert den Fall zur menschlichen Prüfung, bevor eine operative Empfehlung als „präferiert“ gekennzeichnet wird.

Versionierung:

Kostenmodell-Snapshots und Playbooks werden versioniert, damit später nachvollziehbar bleibt, auf welcher Annahmenlage eine Wirkungseinordnung und eine Empfehlung entstanden ist. Das ist insbesondere relevant, wenn sich Kostenannahmen, Prozessparameter oder Standardvorgehen über Zeit ändern. Versionierung ermöglicht damit technische und organisatorische Nachvollziehbarkeit im Fallreview.

Rollen im Agentensystem:

Das Agentensystem ist rollenbasiert, weil unterschiedliche Teilaufgaben unterschiedliche Anforderungen an Datenzugriff, Logik und Begründung haben. Ein Kontextagent rekonstruiert das technische und betriebliche Gesamtbild. Ein Ähnlichkeitsagent durchsucht historische Fälle und Playbooks nach vergleichbaren Mustern. Ein Wirkungsagent übersetzt technische Abweichungen in eine Bandbreite wirtschaftlicher und operativer Folgen. Eine Bandbreite ist bewusst gewählt, weil industrielle Daten unsicher sind und Scheingenauigkeit vermieden werden muss. Ein Maßnahmenagent strukturiert nächste Schritte in Stabilisierung, Diagnose und nachhaltige Behebung. Ein Erkläragent bereitet Ergebnisse zielgruppengerecht auf, einmal für Technik, einmal für Management.

Damit Agenten nicht plausible, aber unbelegte Aussagen erzeugen, ist die Arbeitsweise evidenzgebunden. Evidenzgebunden bedeutet, dass jede Aussage entweder auf referenzierten Messsegmenten, Ereignissen, Kontextdaten oder historischen Fällen basiert, oder als Hypothese gekennzeichnet wird. Hypothese bedeutet hier eine plausible Ursache oder Interpretation, die aber noch nicht durch Daten bestätigt ist.



Im Leitbeispiel bedeutet das: Der Ähnlichkeitsagent liefert referenzierte historische Fälle „Lagerdegradation mit ähnlichem Vibrationsanstieg“, inklusive dokumentierter Maßnahmen und Erfolg. Der Wirkungsagent übersetzt mögliche Stillstandszeit in eine Bandbreite potenzieller Produktionsverluste auf Basis des gültigen Kostenmodell-Snapshots. Der Maßnahmenagent schlägt eine stufenweise Vorgehensweise vor (zunächst Stabilisierung, dann Diagnose, dann nachhaltige Behebung).

Diese Ebene löst die fehlende Priorisierung, weil sie Fallprioritäten aus Kontext, Historie und Wirkung konsolidiert, statt nur Alarmstufen zu wiederholen. Sie löst mangelnde Managementtauglichkeit, weil sie technische Abweichungen in eine verständliche Wirkungseinordnung und Handlungsoptionen übersetzt, inklusive Annahmen und Unsicherheit. Sie löst fehlendes Wissensmanagement, weil sie historische Fälle systematisch nutzbar macht und neue Erkenntnisse strukturiert zurückschreibt.

Ebene 7: Integration und Decision Layer, Nutzen entsteht im Prozess

Ein System ist im Shopfloor nur dann wirksam, wenn es dort erscheint, wo Menschen arbeiten. Das sind typischerweise Leitstandoberflächen, Instandhaltungsprozesse, Ticketing und die Kommunikation zwischen Schichtführung, Instandhaltung und Produktionsleitung.

Wer bekommt wann welche Information:

- Leitstand / Schichtführung erhält die priorisierte Fallliste und eine komprimierte technische Begründung (welche Signale, welcher Verlauf, warum priorisiert), damit schnell entschieden werden kann, welche Fälle sofort Aufmerksamkeit erfordern.
- Instandhaltung erhält die detaillierte technische Fallkarte (Sensorsegmente, Ereigniskette, Datenqualitätskennzeichen, ähnliche Fälle, vorgeschlagene Diagnosepfade), um die technische Prüfung und Maßnahmenplanung vorzubereiten.
- Produktionsleitung erhält eine verdichtete, managementtaugliche Sicht: Wirkung als Bandbreite, Risiko für Liefertermine, sowie Handlungsoptionen und Entscheidungspunkte.

Welche Entscheidungen verpflichtend menschlich sind:

Entscheidungen mit potenzieller Sicherheits- oder Produktionswirkung sind verpflichtend menschlich freizugeben, insbesondere:

- Eingriffe, die den Anlagenbetrieb verändern (zum Beispiel Stop, Bypass, Reduktion der Leistung, Umschaltung)
- Maßnahmen, die Prozessparameter außerhalb definierter Betriebsgrenzen setzen
- Eskalationen, die externe Verpflichtungen betreffen (zum Beispiel Lieferterminentscheidung, Sicherheitsmeldungen)

Feedback-Rückführung:

Rückmeldungen werden nicht informell gesammelt, sondern strukturiert als Teil des Falllebenszyklus erfasst:

- Akzeptiert (Empfehlung umgesetzt)



- Abgelehnt (nicht umgesetzt, mit Begründung)
- Override (alternative Entscheidung, mit Begründung)
- Ergänzend: Outcome-Daten (Zeit bis Stabilisierung, Wiederkehr, Effekt auf Qualität/Ausstoß)

Diese Feedback-Ereignisse werden als Decision Log und Action Outcome in die Lernschicht zurückgeführt. Durch diese strukturierte Rückführung kann der Agent seine Entscheidungslogik kontinuierlich verbessern und zukünftige Empfehlungen auf Basis realer Wirksamkeit adaptiv optimieren.

Wirkungssatz: Diese Ebene stabilisiert die Priorisierung, weil sie Fälle in den operativen Ablauf bringt und Eskalationslogik steuert, statt zusätzliche Listen zu erzeugen. Sie verbessert die managementtaugliche Entscheidungsgrundlage, weil Entscheidungen dokumentiert und begründet im Prozess verankert werden. Sie stärkt Wissensmanagement, weil Entscheidungen und Rückmeldungen als Lernsignal erfasst werden.

Ebene 8: Learning und Knowledge Layer, Lernen wird Teil des Systems

Wissensmanagement in der Industrie scheitert häufig daran, dass Maßnahmen nicht konsequent dokumentiert und Ergebnisse nicht systematisch bewertet werden. Ein Learning-Layer löst das nicht durch eine zusätzliche Wissensablage, sondern durch eine geschlossene Rückkopplung im Falllebenszyklus.

Technisch werden Fall, getroffene Entscheidung, umgesetzte Maßnahmen und Ergebnis miteinander verknüpft. Ergebnis bedeutet hier zum Beispiel, ob die Abweichung verschwunden ist, ob sie wiederkehrte, wie lange die Behebung dauerte oder, ob Nebenwirkungen auf Qualität und Ausstoß entstanden. Daraus entstehen Playbooks, also standardisierte Vorgehensweisen mit Voraussetzungen, Warnhinweisen und Erfolgswahrscheinlichkeiten. Zusätzlich werden Muster formalisiert, damit ähnliche Abweichungen zukünftig schneller erkannt und vergleichbar gemacht werden. Drift Monitoring sorgt dafür, dass Normalverhalten, das sich durch Alterung, Produktwechsel oder Prozessänderungen verschiebt, erkannt wird, damit Modelle und Regeln nachjustiert werden können. Diese Nachjustierung erfolgt kontrolliert über einen spezialisierten Drift-Agenten: Er erkennt Modell- und Regelabweichungen anhand versionierter Baselines, markiert den Drift mit Evidenz (betroffene Merkmale, Betriebsmodi, Zeitfenster) und stößt einen überprüfbaren Aktualisierungsprozess an.

Diese Ebene reduziert langfristig die Warnmeldungsichte ohne klare Priorität, weil das System durch Lernen weniger Rauschen in bearbeitungsrelevante Fälle übersetzt. Sie stärkt die managementtaugliche Entscheidungsgrundlage, weil Wirkungsschätzungen anhand realer Outcomes kalibriert werden. Sie löst fehlendes Wissensmanagement, weil Erfahrungswissen als Playbooks und Musterbibliothek standortübergreifend verfügbar wird.

3.2.2 User story: Schichtverlauf mit AI Agenten

Zu Schichtbeginn im Leitstand ist die zentrale Veränderung nicht ein neues Tool, sondern ein anderes Arbeitsobjekt. Statt vieler gleichwertiger Warnmeldungen sieht die Schichtführung eine kurze Liste priorisierter Fallkarten. Jede Fallkarte bündelt zusammengehörige Alarmer



Messmuster, zeigt den relevanten Zeitverlauf, ordnet den Betriebsmodus ein und verweist auf eine begründete Priorität.

Ein Techniker öffnet die höchste Priorität. Er sieht nicht nur den Alarmtext, sondern ein technisches Gesamtbild: die Sensorsegmente, die das Muster tragen, die dazugehörigen Ereignisse (zum Beispiel Zustandswechsel, Eingriffe, Qualitätshinweise) sowie Hinweise zur Datenqualität. Damit entfällt ein großer Teil des manuellen Zusammenklickens. Die Fallkarte zeigt außerdem, ob es ähnliche Fälle gab und welche Maßnahmen damals wirksam waren. Wenn der erfahrene Kollege nicht verfügbar ist, wird dessen Wissen nicht ersetzt, aber verfügbar gemacht.

Parallel bekommt die Produktionsleitung eine verdichtete Sicht. Sie sieht keine Rohkennzahlen, sondern eine Wirkungseinordnung als Bandbreite, inklusive der Annahmen, auf denen sie basiert. So wird aus einer technischen Abweichung eine nachvollziehbare Entscheidungsgrundlage, ohne dass operative Experten diese erst in Meetings und Tabellen übersetzen müssen.

Die nächsten Schritte sind als abgestufter Maßnahmenplan dargestellt. Zunächst Stabilisierung, um das Risiko zu begrenzen, dann Diagnose, um Ursachenhypothesen zu prüfen, und schließlich nachhaltige Behebung, um Wiederkehr zu verhindern. Wenn eine Empfehlung einen Eingriff in den Betrieb bedeutet, ist die Freigabe explizit erforderlich. Damit wird Geschwindigkeit nicht auf Kosten von Sicherheit erkaufte.

Nach Umsetzung wird die Maßnahme nicht nur abgeschlossen, sondern bewertet. Der Fall wird gelöst, Entscheidung und Ergebnis werden verknüpft, und wenn ein Muster erfolgreich behoben wurde, entsteht daraus ein Playbook. Das System wird über Zeit konsistenter, weil Lernen nicht als Extra-Aufgabe, sondern als integrierter Schritt im Fallprozess stattfindet.

3.2.3 Der End-to-End-Datenfluss im IoT-Abweichungsmanagement

Datenfluss und Prompt-Blueprints im IoT-Abweichungsmanagement: Warum ein expliziter Datenfluss notwendig ist

In IoT-gestützten Industrieumgebungen entsteht Komplexität nur selten durch einzelne Analyseverfahren oder Modelle. Sie entsteht vor allem an den Übergängen: zwischen Sensorik und Analyse, zwischen technischer Bewertung und organisatorischer Bearbeitung sowie zwischen Entscheidung und Lernen. Ohne einen explizit definierten Datenfluss entstehen Brüche, etwa wenn ein Alarm zwar erkannt, aber ohne Kontext bewertet wird, oder wenn Maßnahmen umgesetzt werden, ohne dass ihr Ergebnis systematisch zurückgeführt wird.

Dieses Ergänzungskapitel beschreibt den Datenfluss deshalb nicht abstrakt, sondern als Lebenszyklus eines Abweichungsfalls auf dem Shopfloor. Es zeigt, wie technische Daten schrittweise in bearbeitbare Fallkarten überführt werden und wie spezialisierte AI-Agenten diesen Prozess mit evidenzgebundener Entscheidungsunterstützung begleiten. Ziel ist es, Nachvollziehbarkeit, Konsistenz und Lernfähigkeit technisch abzusichern.

Der Datenfluss als Lebenszyklus eines Falls:

Ein Abweichungsfall entsteht nicht in einem einzelnen Schritt. Er durchläuft mehrere klar voneinander getrennte Phasen, die jeweils definieren, welche Daten vorliegen, wie sie interpretiert werden dürfen und welche Entscheidungen zulässig sind.



Zu Beginn werden Sensorwerte und Ereignisse aufgenommen und qualitätsgesichert. Darauf aufbauend werden Merkmale und Kontextinformationen erzeugt, die technische Muster beschreiben. Erst danach identifiziert die Abweichungserkennung auffällige Konstellationen und formuliert diese als Kandidaten. Wichtig ist hierbei die bewusste Abgrenzung: Die Abweichungserkennung erzeugt keine Diagnose, sondern eine formal begründete Auffälligkeits-Hypothese, deren Evidenz explizit im Kandidatenobjekt gespeichert wird.

In einem nächsten Schritt bündelt das Case Management mehrere Kandidaten zu einem organisatorisch steuerbaren Fall. Dieser Fall wird anschließend durch die Agentenebene angereichert, sodass Priorität, Wirkung, Maßnahmen und Begründung entstehen. Die daraus resultierende Entscheidung wird in operative Systeme zurückgeführt und abschließend als Lernergebnis wieder in das System eingespeist.

Dieser Lebenszyklus ist die technische Grundlage dafür, dass aus Sensorrauschen eine nachvollziehbare Entscheidungsunterstützung entsteht.

Zentrale Übergänge im Datenfluss:

Entlang dieses Lebenszyklus lassen sich einige zentrale Übergänge identifizieren, die für den Shopfloor besonders relevant sind. Aus Sensorwerten und Ereignissen entstehen zunächst Merkmalsdaten, etwa Trends, Sprungstellen oder Fensterstatistiken, ergänzt um den aktuellen Betriebsmodus. Diese Merkmale machen Muster sichtbar und interpretierbar.

Aus den Merkmalsdaten entstehen anschließend Kandidaten mit einer Bewertung, zugehörigen Evidenzsegmenten und einer Unsicherheitsmarkierung. Der Zweck ist es, Auffälligkeiten zu kennzeichnen, ohne sofort operative Arbeit auszulösen.

Mehrere Kandidaten werden im Case Management zu einer Fallkarte gebündelt, die Status, Verantwortlichkeit, Fristen und zusammengeführte Evidenzen enthält. Erst diese Fallkarte ist das Arbeitsobjekt für Leitstand, Instandhaltung und Produktionsleitung.

Auf Basis der Fallkarte werden Wirkungseinordnung und Maßnahmenpläne erzeugt. Die Wirkung wird als Bandbreite beschrieben, nicht als exakter Wert, um Unsicherheit transparent zu halten. Maßnahmen werden strukturiert vorgeschlagen, ohne automatische Eingriffe auszulösen.

Abgeschlossene Entscheidungen und ihre Ergebnisse fließen schließlich in Wissensobjekte wie Playbooks, Musterbibliotheken und Erfolgswahrscheinlichkeiten zurück. Dadurch wird Erfahrungswissen reproduzierbar und skalierbar.

Zentrale Datenobjekte im Überblick:

Ein Sensor Reading bezeichnet einen einzelnen Messpunkt mit Zeitstempel, Messwert und Qualitätskennzeichnung. Ein Sensor Segment fasst viele solcher Messpunkte in einem Zeitfenster zusammen, weil technische Muster in der Praxis über Zeit entstehen. Ein Event Record beschreibt ein Ereignis wie einen Alarm, einen Zustandswechsel oder eine Qualitätsmeldung, die häufig Kontext oder Ursache liefert.

Das Asset Master Objekt beschreibt die betroffene Produktionslinie, Maschine, das Aggregat oder Netzsegment inklusive Kritikalität und Hierarchie. Der Operational Context ergänzt diese Informationen um den betrieblichen Rahmen, etwa Schicht, Auftrag oder Betriebsmodus.



Ein Deviation Candidate ist die technische Repräsentation einer Auffälligkeits-Hypothese. Ein Deviation Case ist das organisatorische Arbeitsobjekt mit Status, Verantwortlichkeit und Frist. Similar Case Sets bündeln vergleichbare historische Fälle. Ein Cost Model Snapshot hält versionierte Kostenannahmen fest, um spätere Wirkungsabschätzungen nachvollziehbar zu machen. Impact Assessments, Recommendation Plans und Explanation Packets bilden schließlich Wirkung, Maßnahmen und Begründung strukturiert ab. Decision Logs und Action Outcomes erfassen Entscheidungen und Ergebnisse als Lernsignal.

Agenteninteraktion entlang des Falllebenszyklus:

Die Agenteninteraktion wird zentral durch einen Orchestrator gesteuert. Diese Rolle ist notwendig, um parallele, widersprüchliche oder unvollständige Analysen zu vermeiden. Der Orchestrator definiert die Reihenfolge der Agentenaktivierung, prüft die Vollständigkeit der Ergebnisse und stellt sicher, dass alle Ausgaben einem definierten Format folgen.

Im Standardablauf beginnt die Agentenarbeit nach der Erstellung einer Fallkarte. Zunächst wird durch den Context Agent der operative Gesamtzusammenhang rekonstruiert. Anschließend identifiziert der Similarity Agent vergleichbare historische Fälle. Darauf aufbauend schätzt der Impact Agent die potenzielle Wirkung ab. Der Recommendation Agent strukturiert daraufhin mögliche Maßnahmen, und der Explanation Agent bereitet die Ergebnisse verständlich auf.

Ohne Kontext bleibt Ähnlichkeit unscharf, ohne Ähnlichkeit bleiben Maßnahmen generisch, ohne Wirkung fehlt Priorisierung, und ohne Erklärung sinkt die Akzeptanz der Entscheidung.

Input- und Output der Agenten:

Alle Agenten arbeiten ausschließlich auf freigegebenen, versionierten Datenobjekten. Dazu zählen insbesondere Fallkarten, Sensorsegmente, Ereignisketten, Cost-Model-Snapshots und Playbooks. Freitext oder implizite Annahmen außerhalb dieser Objekte sind nicht zulässig.

Die Ergebnisse der Agenten werden strukturiert abgelegt. Jede Analyse liefert explizit ausgewiesene Prioritäten, Wirkungsbreiten, empfohlene Maßnahmen, Begründungen sowie offene Punkte oder Unsicherheiten. Dadurch bleibt der Entscheidungsprozess auditierbar und überprüfbar.

Konfliktauflösung und menschliche Entscheidungsinstanzen:

Wenn Agenten zu widersprüchlichen Ergebnissen kommen, etwa wenn historische Fälle eine geringe Kritikalität nahelegen, während die Wirkungsschätzung hoch ausfällt, markiert der Orchestrator diesen Konflikt explizit. In solchen Fällen fordert er zusätzliche Daten an oder eskaliert den Fall zur menschlichen Prüfung.

Grundsätzlich gilt: Alle Eingriffe mit potenzieller Auswirkung auf Sicherheit, Anlagenzustand oder Produktionsoutput erfordern eine explizite menschliche Freigabe. Das System unterstützt Entscheidungen, es ersetzt sie nicht.

Feedback und Versionierung als Grundlage für Lernen:

Nach der Entscheidung wird Feedback technisch erfasst. Operative Nutzer können Empfehlungen akzeptieren, ablehnen oder überschreiben und müssen dabei eine Begründung



angeben. Dieses Feedback wird gemeinsam mit dem tatsächlichen Ergebnis des Falls gespeichert.

Kostenmodelle und Playbooks werden versioniert, damit nachvollziehbar bleibt, auf welcher Annahmenbasis eine Entscheidung getroffen wurde. Diese Versionierung ist essenziell, um Entscheidungen auch retrospektiv bewerten und vergleichen zu können.

Durch diese Rückkopplung entsteht ein lernendes System, das nicht nur Abweichungen erkennt, sondern seine Entscheidungsqualität über Zeit systematisch verbessert.

Fazit

Dieses Kapitel zeigt, dass der eigentliche Engpass im industriellen Abweichungsmanagement nicht in fehlender Sensorik oder Analyseverfahren liegt, sondern in der Übersetzung: von Rohsignalen zu priorisierten, begründbaren und handlungsfähigen Entscheidungen. In asset-intensiven Umgebungen entstehen relevante Risiken selten durch einzelne Alarme, sondern durch Muster über Zeit, Kontext und mehrere Signale hinweg. Genau hier setzt die dargestellte Referenzarchitektur an.

Der zentrale Mehrwert der Architektur liegt nicht in zusätzlicher „Intelligenz“, sondern in Struktur. Durch die konsequente Trennung von Datenaufnahme, Musterbildung, Abweichungserkennung, Fallsteuerung und Entscheidungsunterstützung wird verhindert, dass Alarmflut technisch verstärkt oder operative Verantwortung unklar verteilt wird. Stattdessen entsteht ein durchgängiger Fallebenszyklus, in dem jede Ebene einen klaren Beitrag zur Reduktion von Rauschen, zur Priorisierung von Arbeit und zur Nachvollziehbarkeit von Entscheidungen leistet.

Besonders entscheidend ist der Perspektivwechsel vom Alarm zur Fallkarte. Die Fallkarte ist das gemeinsame Arbeitsobjekt von Leitstand, Instandhaltung und Management. Sie bündelt technische Evidenz, betrieblichen Kontext, eine begründete Priorität, eine transparente Wirkungseinordnung sowie abgestufte Maßnahmen. Damit wird aus einem technischen Signal ein steuerbarer Vorgang und aus implizitem Erfahrungswissen reproduzierbare Entscheidungslogik.

Die Integration von AI-Agenten verstärkt diesen Effekt, ohne operative Verantwortung zu automatisieren. Die Agenten fungieren nicht als Black Box, sondern als strukturierte Entscheidungsunterstützung: evidenzgebunden, versioniert und konfliktfähig. Sie machen historische Erfahrung systematisch nutzbar, übersetzen technische Abweichungen in managementtaugliche Wirkungsbandbreiten und erhöhen die Konsistenz von Entscheidungen über Schichten, Standorte und Personen hinweg. Gleichzeitig bleibt die menschliche Freigabe dort verpflichtend, wo Sicherheit, Anlagenzustand oder Produktionsleistung betroffen sind.

Langfristig entfaltet die Architektur ihre volle Wirkung im Learning-Layer. Erst durch die geschlossene Rückkopplung von Entscheidung, Maßnahme und Ergebnis entsteht ein System, das nicht nur Abweichungen erkennt, sondern seine Entscheidungsqualität kontinuierlich verbessert. Playbooks, Musterbibliotheken und kalibrierte Wirkungsannahmen machen Lernen zum integralen Bestandteil des operativen Prozesses, nicht zu einer zusätzlichen Dokumentationsaufgabe.



Dualer

CONSULTING

Zusammenfassend zeigt dieses Kapitel: Wirksames IoT-gestütztes Abweichungsmanagement entsteht durch eine saubere End-to-End-Architektur, die technische Realität, operative Abläufe und Managementanforderungen konsequent zusammenführt. Dort, wo aus Sensorrauschen wenige, begründete Fallkarten werden, entsteht echte Handlungsfähigkeit auf dem Shopfloor.



3.3 Rechtliche Umsetzung

Dieses Kapitel analysiert die rechtlichen Risiken der beschriebenen IoT- und AI-gestützten Abweichungsmanagement-Architektur entlang ihrer technischen Ebenen. Ziel ist nicht eine vollständige Rechtsprüfung, vielmehr die Identifikation der jeweils signifikantesten rechtlichen Konfliktlinie je Ebene sowie die Ableitung konkreter rechtlicher Umsetzungsmaßnahmen. Maßgeblich sind insbesondere der EU AI Act (VO (EU) 2024/1689), die DSGVO, haftungsrechtliche Grundsätze sowie arbeitsrechtliche Schutzmechanismen.

Auf **Ebene 1 (Sensorik und Edge-Verarbeitung)** entsteht das zentrale rechtliche Risiko aus der haftungsrechtlichen Relevanz fehlerhafter oder unzureichend gekennzeichnete Messdaten. Auch wenn auf dieser Ebene noch keine AI-gestützte Entscheidungslogik greift, bilden die hier erzeugten Daten die Grundlage für spätere Priorisierungen und Handlungsempfehlungen. Fehlerhafte oder unvollständig gekennzeichnete Sensordaten fließen in später generierte Fallkarten ein, die als Entscheidungsgrundlage dienen.

Um eine haftungsrechtliche Zurechnung fehlerhafter Entscheidungen zu begrenzen, muss die Architektur sicherstellen, dass jede Weiterverarbeitung an eine dokumentierte Datenqualitätskennzeichnung gebunden ist. Zudem müssen nicht verlässliche Messwerte systematisch von wirkungsrelevanten Bewertungen ausgeschlossen werden (beispielsweise durch eine Nicht-Eignungsklausel). Die Festlegung von Datenqualitätskriterien sowie die Einstufung der Messwerte erfolgt durch den Systembetreiber auf Ebene der Edge-Komponenten und wird als verpflichtendes Metadatum jedem Messwert beigefügt. Die Weiterverarbeitung auf nachgelagerten Ebenen ist technisch nur für Messwerte zulässig, die als für entscheidungsunterstützende Zwecke geeignet gekennzeichnet sind. Nicht geeignete Messwerte bleiben auf eine rein dokumentarische Nutzung beschränkt.

Auf **Ebene 2 (Data Ingestion)** liegt das maßgebliche Risiko im Bereich der datenschutzrechtlichen Zweckbindung. Die Integration externer Dokumente wie Wartungsprotokolle oder Wissensdatenbanken (aus Confluence Seiten oder SharePoint-Dokumenten) kann personenbezogene Informationen enthalten. Wird dieses Material ohne klare Zweckdefinition in ein AI-gestütztes System überführt, besteht das Risiko, dass Daten, die ursprünglich für Dokumentation oder Kommunikation gedacht waren, für Entscheidungs- oder Lernprozesse zweckentfremdet werden.

Um eine unzulässige Zweckausweitung im Sinne von Art. 5 DSGVO zu vermeiden, ist eine klare Trennung zwischen technischem Wissen und personenbezogenen Metadaten erforderlich. Zu ergänzen ist dies durch Pseudonymisierung und eine explizite Zweckdefinition (technische Wissensunterstützung, nicht Leistungsbewertung) der eingebundenen Inhalte. Die Zweckbindung der eingebundenen Inhalte wird durch den verantwortlichen Fachbereich bereits beim Import der Daten festgelegt und in Form technischer Metadaten im Ingestion-Layer verankert. Personenbezogene Inhalte werden automatisiert pseudonymisiert oder vom inhaltlichen Wissenskern getrennt, sodass sie ausschließlich für die definierte Wissensunterstützung genutzt werden können. Eine Nutzung zur Leistungs- oder Verhaltensbewertung ist systemseitig ausgeschlossen.

Auf **Ebene 3 (Datenverarbeitung und Feature Engineering)** entsteht ein Transparenzrisiko, da abstrakte Merkmale (Trends, Varianz, Korrelationen) und Muster gebildet werden, die für



menschliche Nutzer nicht unmittelbar nachvollziehbar sind, aber die spätere Bewertung von Abweichungen maßgeblich beeinflussen. Ohne Dokumentation entsteht ein Black-Box-Vorwurf, insbesondere wenn Entscheidungen später wirtschaftliche oder organisatorische Konsequenzen haben.

Rechtlich kritisch wird dies, wenn diese Merkmale faktisch als Vorstufe einer automatisierten Bewertung wirken (Art. 13, 14 DSGVO; Art. 12 EU AI Act). Die rechtliche Umsetzung erfordert daher eine dokumentierte und nachvollziehbare Beschreibung der verwendeten Merkmale sowie eine klare Abgrenzung zwischen rein beschreibenden Features und bewertenden Schlussfolgerungen.

Die Definition und Beschreibung der gebildeten Merkmale erfolgt zentral durch das Entwicklungsteam und wird in einem fortlaufend gepflegten Feature-Katalog dokumentiert. Dieser enthält Zweck, Herkunft und Wirkungsbereich jedes Merkmals und grenzt ausdrücklich beschreibende Aggregationen von bewertenden Ableitungen ab. Der Katalog dient sowohl der internen Nachvollziehbarkeit als auch der externen Prüfbarkeit.

Auf **Ebene 4 (Abweichungserkennung)** liegt das Risiko in einer verdeckten Vorentscheidung durch Scoring- oder Klassifikationsmechanismen. Auch wenn formell keine Entscheidung getroffen wird, kann die Markierung bestimmter Ereignisse als relevant die menschliche Aufmerksamkeit lenken und andere Fälle faktisch ausblenden. Damit entsteht eine entscheidungslenkende Wirkung, auch ohne formale Entscheidung.

Um den Anforderungen an wirksame menschliche Aufsicht nach dem Art. 14 des EU AI Acts zu genügen, muss die Architektur Abweichungen ausdrücklich als Hypothesen kennzeichnen und eine automatische Überführung in Maßnahmen oder Prioritäten ausschließen. Die Kennzeichnung von Abweichungen als Hypothesen wird durch das System selbst vorgenommen und in der Nutzeroberfläche eindeutig sprachlich und visuell kenntlich gemacht. Eine automatische Priorisierung oder Maßnahmeneinleitung findet auf dieser Ebene nicht statt, sodass die Bewertung der Relevanz ausschließlich durch den menschlichen Nutzer erfolgt. Dadurch bleibt die Entscheidungsverantwortung vollständig außerhalb des Systems.

Auf **Ebene 5 (Case Management)** entsteht ein datenschutz- und arbeitsrechtliches Risiko durch den impliziten Personenbezug von Fallzuweisungen, Bearbeitungsstatus und Zeitverläufen. Fallkarten enthalten Verantwortlichkeiten, Bearbeitungsstatus und Zeitverläufe. Über längere Zeiträume können daraus Leistungs- oder Verhaltensprofile entstehen.

Zur rechtlichen Absicherung ist eine rollenbasierte statt personenbezogener Zuordnung erforderlich. Zu ergänzen ist dies durch Lösch- und Anonymisierungskonzepte für abgeschlossene Fälle sowie eine klare Trennung zwischen operativer Dokumentation und personalbezogener Bewertung. Mit diesen Maßnahmen wird dem Schutz vor unzulässiger Leistungsüberwachung nach Art. 6, Art. 88 DSGVO; § 26 BDSG genüge getan. Die Zuordnung von Fällen erfolgt organisatorisch über Rollen und Funktionen und wird technisch ohne dauerhafte personenbezogene Identifikatoren umgesetzt. Verantwortlichkeiten und Bearbeitungsstände werden nach Abschluss eines Falls gemäß einem definierten Lösch- und Anonymisierungskonzept bereinigt. Eine Nutzung der Fallhistorie für personalbezogene Auswertungen ist weder vorgesehen noch systemisch möglich.

Die **Ebene 6 (KI-Agenten)** stellt den zentralen rechtlichen Brennpunkt dar. Hier werden technische Analysen in priorisierte Handlungsempfehlungen mit wirtschaftlicher oder



sicherheitsrelevanter Wirkung übersetzt. Auch ohne formale Automatisierung besteht die Gefahr einer faktischen Entscheidungslenkung, wenn Empfehlungen regelmäßig übernommen werden.

Um die Anforderungen des EU AI Act (insbesondere Art. 6, Art. 9, Art. 14 EU AI Act) an Hochrisiko-Systeme zu erfüllen, muss die Architektur sicherstellen, dass keine Default-Optionen existieren, Unsicherheiten sichtbar gemacht werden und jede Übernahme einer Empfehlung eine bewusste menschliche Entscheidung voraussetzt. Handlungsempfehlungen der AI-Agenten werden ausschließlich in einer nicht vorstrukturierten Form bereitgestellt und enthalten keine vorausgewählten Optionen oder impliziten Entscheidungsvorschläge. Unsicherheiten, Annahmen und alternative Handlungsoptionen werden explizit ausgewiesen, sodass jede Übernahme einer Empfehlung eine bewusste menschliche Entscheidung erfordert. Die Verantwortung für die Umsetzung verbleibt jederzeit beim Nutzer.

Auf **Ebene 7 (Decision Logs und Systemintegration)** liegt das größte DSGVO-Risiko in der langfristigen Speicherung von Entscheidungsverläufen in Verbindung mit Personen oder Rollen. In Kombination mit Outcomes können so personenbezogene Erfolgskennzahlen entstehen.

Die rechtliche Umsetzung (unter Betracht von Art. 5 Abs. 1 lit. c, e DSGVO) erfordert daher eine strikte Datenminimierung, eine zeitliche Begrenzung der Speicherung sowie die Entkopplung von Lernsignalen und personenbezogenen Identitäten, um eine Zweckentfremdung zu verhindern. Entscheidungsverläufe werden nur in dem Umfang gespeichert, der für Nachvollziehbarkeit und Systemstabilität erforderlich ist, und unterliegen festen zeitlichen Begrenzungen. Personen- oder rollenbezogene Bezüge werden frühzeitig entfernt oder aggregiert, bevor die Daten für Analyse- oder Lernzwecke herangezogen werden. Eine Verknüpfung mit individuellen Leistungskennzahlen ist ausgeschlossen.

Auf **Ebene 8 (Learning- und Knowledge-Layer)** besteht schließlich das Risiko der Verstetigung fehlerhafter oder kontextabhängiger Annahmen über Zeit. Historische Entscheidungen prägen zukünftige Empfehlungen und können bei veränderten Rahmenbedingungen zu systematischen Fehlbewertungen führen.

Um den Anforderungen an Risikomanagement und Robustheit (Art. 9, Art. 15 EU AI Act) nach dem EU AI Act gerecht zu werden, muss Lernen stets versioniert, überwacht und an eine explizite menschliche Validierung gebunden sein. Lernprozesse werden versioniert und durch den verantwortlichen Fach- oder Systemverantwortlichen regelmäßig überprüft und freigegeben. Änderungen am Wissensbestand erfolgen nur nach dokumentierter Validierung und sind jederzeit auf ihren Entstehungskontext zurückführbar. Dadurch wird sichergestellt, dass historische Annahmen nicht unreflektiert in veränderte Entscheidungskontexte übernommen werden.

Zusammenfassung und Bewertung

Die rechtlichen Hauptrisiken der Architektur entstehen nicht in der Sensorik, sondern dort, wo technische Erkenntnisse in strukturierte Entscheidungslogik und langfristiges Wissen überführt werden. Entscheidend für die Rechtskonformität ist weniger die Existenz technischer Kontrollmechanismen als deren konsequente organisatorische Durchsetzung im Betrieb.



3.4 Abschließendes Fazit zum Use Case IoT Abweichungsmanagement

Das IoT-gestützte Abweichungsmanagement entfaltet seinen Mehrwert nicht durch zusätzliche Alarme, sondern durch die strukturierte Überführung von Rohsignalen in priorisierte, begründete Fallkarten. Durch die Trennung von Datenaufnahme, Musterbildung, Fallsteuerung und evidenzgebundener Entscheidungsunterstützung werden Alarmflut, mangelnde Managementtauglichkeit und fehlendes Wissensmanagement systematisch adressiert. So entsteht eine nachvollziehbare, skalierbare Entscheidungslogik, die operative Handlungsfähigkeit stärkt, ohne menschliche Verantwortung zu ersetzen.



4. Use Case 2 | Management Decision Agent

In Management-Kontexten müssen regelmäßig strategische und operative Entscheidungen unter hohem Zeitdruck getroffen werden. Gleichzeitig sind die dafür relevanten Informationen häufig über verschiedene Systeme verteilt, unterschiedlich aufbereitet und nicht konsistent miteinander verknüpft.

Die folgenden Pain Points beschreiben zentrale Herausforderungen in der Entscheidungsfindung und zeigen, wo ein Management Decision Agent ansetzt, um Transparenz, Struktur und eine verlässliche Entscheidungsgrundlage zu schaffen.

Fragmentierte Entscheidungsgrundlagen werden konsolidiert - für proaktives, transparentes Management

Use Case: Management Decision Agent



4.1 Pain Points

Pain Point 1: Fragmentierte und inkonsistente Entscheidungsgrundlagen

Management-Entscheidungen wie Budgetfreigaben, Projektpriorisierungen oder Ressourcenallokationen erfordern eine ganzheitliche, konsistente und aktuelle Sicht auf die Unternehmenssituation. In der Praxis liegen die hierfür notwendigen Informationen jedoch verteilt über unterschiedliche Systeme, Reports und Präsentationsformate vor. Finanzkennzahlen, Projektfortschritte, Kapazitätsauslastungen und Risikoeinschätzungen folgen häufig unterschiedlichen Logiken, Aktualisierungszyklen und Interpretationen. Eine einheitliche Entscheidungsgrundlage existiert daher in vielen Fällen nicht.

Entscheidungsrelevante Informationen werden vor Management-Terminen manuell zusammengeführt, oftmals unter hohem Zeitdruck. Abweichungen und Unstimmigkeiten werden situativ diskutiert, ohne systematisch aufgelöst zu werden. Für das Management bedeutet dies eine erhöhte Unsicherheit bei Entscheidungen sowie eine starke Abhängigkeit von einzelnen Personen, die Daten interpretieren und einordnen. Ein erheblicher Teil der verfügbaren Zeit wird für die Klärung von Zahlen aufgewendet, anstatt strategische Handlungsoptionen zu bewerten.

Der Management Decision Agent adressiert dieses Problem, indem er entscheidungsrelevante Daten automatisiert aus bestehenden Systemen integriert, konsolidiert und transparent



aufbereitet. Dadurch entsteht eine belastbare und konsistente Entscheidungsgrundlage, die Diskussionen auf inhaltliche Fragestellungen lenkt und die Qualität von Management-Entscheidungen nachhaltig erhöht.

Pain Point 2: Zeitverzögerte und reaktive Entscheidungsfindung

In vielen Organisationen erfolgt die Management-Steuerung überwiegend auf Basis vergangenheitsorientierter Reports. Kritische Entwicklungen werden häufig erst erkannt, wenn Abweichungen bei Budget, Zeit oder Qualität bereits eingetreten sind und der Handlungsspielraum deutlich eingeschränkt ist. Entscheidungsgrundlagen werden in festen Zyklen erstellt und bilden die tatsächliche Dynamik laufender Projekte oder Initiativen nur unzureichend ab.

Frühindikatoren für Risiken oder Zielabweichungen sind in den vorhandenen Daten zwar enthalten, werden jedoch nicht kontinuierlich analysiert, priorisiert oder in einen Entscheidungszusammenhang gebracht. Entscheidungen erfolgen daher häufig unter Zeitdruck und sind von Eskalationen sowie kurzfristigen Korrekturmaßnahmen geprägt. Für Führungskräfte entsteht der Eindruck eines reaktiven Steuerungsmodus, in dem permanentes Eingreifen notwendig ist, anstatt gezielt und vorausschauend agieren zu können.

Der Management Decision Agent unterstützt das Management, indem er relevante Kennzahlen kontinuierlich überwacht, Trends und Muster frühzeitig identifiziert und potenzielle Risiken transparent macht. Entscheidungen können dadurch proaktiver, planbarer und mit größerem strategischem Handlungsspielraum getroffen werden.

Pain Point 3: Hohe kognitive Belastung und personenabhängige Entscheidungen

Management-Entscheidungen sind zunehmend komplex und erfordern die gleichzeitige Berücksichtigung einer Vielzahl von Einflussfaktoren. In der Praxis fehlen jedoch häufig strukturierte Entscheidungslogiken, sodass Informationen implizit gewichtet und auf Basis individueller Erfahrungen interpretiert werden. Handlungsoptionen werden in Meetings diskutiert, ohne systematisch gegenübergestellt oder hinsichtlich ihrer Auswirkungen transparent bewertet zu werden.

Dies führt zu einer hohen kognitiven Belastung für Führungskräfte sowie zu einer starken Abhängigkeit von einzelnen Meinungsführern oder Wissensträgern. Entscheidungen sind dadurch weniger konsistent, schwerer nachvollziehbar und im Nachhinein nur eingeschränkt erklärbar.

Der Management Decision Agent schafft hier Abhilfe, indem er Entscheidungsoptionen strukturiert aufbereitet, Annahmen explizit macht und Auswirkungen vergleichbar darstellt. Er reduziert Komplexität, ohne Entscheidungen zu automatisieren, und unterstützt das Management dabei, fundierte, nachvollziehbare und konsistente Entscheidungen zu treffen.



4.2 Technische Umsetzung

Ausgangsszenario & Zielsetzung

Management-Entscheidungen (z. B. Budgetfreigaben, Projektpriorisierung, Ressourcenallokation oder Eskalationen bei Abweichungen) scheitern in der Praxis selten an fehlenden Daten, sondern daran, dass die Informationen verteilt, unterschiedlich definiert und zeitlich nicht synchron vorliegen. Dies liegt unter anderem daran, dass Daten mit unterschiedlichen Zeitstempeln verarbeitet werden (z. B. Buchungsdatum vs. Reporting-Datum), auf verschiedenen Ebenen aggregiert sind (Projekt- vs. Arbeitspaketebene) und keine einheitliche, zentral definierte KPI-Logik existiert. Entscheidungsunterlagen entstehen deshalb häufig durch manuelle Konsolidierung, sind schwer vergleichbar und hängen stark von einzelnen Personen ab.

Das Zielbild des Management Decision Agent ist ein technischer Aufbau, der diese Lücke schließt: Der Agent sammelt relevante Daten automatisiert aus bestehenden Systemen, vereinheitlicht KPI-Definitionen und stellt dem Management „decision-ready“ Inhalte bereit, als standardisiertes Decision Pack für Meetings sowie als Frühwarnsystem im Tagesgeschäft. Wichtig ist dabei der Dualer-Ansatz: Der Agent unterstützt Entscheidungen, automatisiert sie aber nicht. Die Entscheidungshoheit bleibt vollständig beim Management.

Einführung des Leitbeispiels

Als Leitbeispiel dient ein typisches Portfolio-Szenario: Ein Unternehmen steuert 10-20 Initiativen oder Projekte, die unterschiedliche Fachbereiche betreffen und um Budgets sowie Kapazitäten konkurrieren. Das Management trifft in regelmäßigen Steuerungsterminen Entscheidungen über Prioritäten, Budgetanpassungen, Ressourcenumverteilung oder Stop/Delay von Vorhaben.

In der aktuellen Realität liegen die dafür relevanten Informationen verteilt in ERP/Finanzsystemen, Projekttools, BI-Reports und Excel-Dateien. Der Management Decision Agent setzt genau dort an und erzeugt aus diesen Quellen eine konsistente Management-Sicht: Was ist der Status? Wo sind Abweichungen? Welche Risiken sind kritisch? Was hat sich seit dem letzten Termin verändert? Welche Optionen gibt es und welche Trade-offs entstehen?

4.2.1 Architektur

Datenquellen

Die technische Umsetzung startet bewusst pragmatisch: Es werden zunächst nur 2-3 Kernquellen angebunden, um schnell einen Minimum Viable Product (MVP) aufzubauen. Typische Startquellen sind:

- ein Finanz-/ERP-System (Budget, Ist-Kosten, Forecast)
- ein Projektmanagement-Tool (Status, Meilensteine, Risiken, Abhängigkeiten)
- ein BI- oder Excel-Reporting (Aggregationen, historisierte Reports)

Je nach Organisation können weitere Quellen wie:

- Ticketing-Systeme (ServiceNow)
- Dokumentationsplattformen (Confluence/SharePoint) oder
- Data-Warehouse-Lösungen (Snowflake, BigQuery)

ergänzt werden.



Ein zentraler Grundsatz für den Piloten ist „read-only“: Der Agent erhält zunächst nur Lesezugriff. Dadurch ist die Einführung risikoarm, bestehende Prozesse werden nicht gestört, und die Organisation kann den Mehrwert in kurzer Zeit validieren.

Datenverarbeitung (Layer-Modell)

Damit die Ergebnisse managementtauglich sind, braucht es eine klare Trennung von Verantwortlichkeiten in der Datenverarbeitung. In der Praxis hat sich ein Layer-Modell bewährt, das technisch wie fachlich sauber aufeinander aufbaut:

- **Connector & Ingestion Layer:** In diesem Layer werden Daten aus den Quellsystemen über APIs, Exporte oder Standard-Connectoren angebunden. Ziel ist eine stabile, automatisierte Datenversorgung mit definierten Update-Zyklen (z. B. täglich oder mehrmals wöchentlich).
- **Data Processing Layer (Harmonisierung):** In diesem Layer werden Rohdaten technisch bereinigt und vereinheitlicht. Dazu zählen unter anderem die Entfernung von Duplikaten, der Umgang mit fehlenden Werten, Format- und Typkonvertierungen sowie grundlegende Validierungen (z. B. Wertebereiche, Pflichtfelder). Zusätzlich werden technische Qualitätsregeln etabliert, etwa zur Datenvollständigkeit, Aktualität (Data Freshness) und strukturellen Plausibilität. Dieser Layer beantwortet bewusst nicht die Frage, was eine Kennzahl fachlich bedeutet, sondern stellt sicher, dass die Daten technisch korrekt, konsistent und verlässlich für die nächste Verarbeitungsebene vorliegen.
- **Semantic & KPI Layer:** In diesem Layer erfolgt die fachliche Harmonisierung der Daten. KPI-Definitionen werden zentral festgelegt, beispielsweise was unter „Budget“, „Forecast“, „Fortschritt“ oder „Risiko“ konkret zu verstehen ist und wie diese Kennzahlen berechnet werden. Hier werden außerdem Geschäftslogiken, Schwellenwerte und Ampellogiken definiert (z. B. Budgetabweichung > X % = rot). Die fachliche Verantwortung für diese Definitionen liegt initial beim Management bzw. bei den zuständigen Fachbereichen. Der AI-Agent übernimmt diese Definitionen nicht autonom, sondern nutzt sie konsistent, dokumentiert und reproduzierbar. Dies hat den großen Vorteil, dass Zahlen vergleichbarer werden und Diskussionen drehen sich nicht mehr um Definitionen, sondern um Entscheidungen.
- **Decision-Ready Layer:** Der letzte Layer bereitet die Informationen so auf, wie sie im Management wirklich genutzt werden: als standardisierte Decision Packs (Status, Top-Abweichungen, Risiken, Abhängigkeiten, Veränderungen seit dem letzten Termin) sowie als priorisierte Executive Views (z. B. „Top 5 Risiken“, „größte Budgetabweichungen“, „kritische Pfade“). Der AI-Agent kann auf dieser Basis erste Handlungsempfehlungen oder Entscheidungsoptionen vorschlagen, ohne selbst Entscheidungen zu treffen oder operative Maßnahmen auszulösen. Jede Kennzahl ist mit Quelle, Zeitstempel und zugrunde liegender Definition versehen, um Transparenz und Nachvollziehbarkeit sicherzustellen. Rückkopplungen in vorgelagerte Layer erfolgen ausschließlich kontrolliert über einen Orchestrator, der entsprechende Folgeaktionen steuert und freigibt.

Dieses Layer-Modell ist nicht „theoretisch“, sondern ein praktischer Bauplan: Es verhindert, dass ein AI-Agent direkt auf ungeprüfte Rohdaten losgelassen wird und schafft die Grundlage für verlässliche, reproduzierbare Entscheidungsunterstützung.



Integration von AI-Agenten

Auf Basis der strukturierten Layer können AI-Agenten gezielt dort eingesetzt werden, wo sie den höchsten Mehrwert liefern: bei Analyse, Priorisierung, Zusammenfassung und Entscheidungsunterstützung.

Ein bewährter Ansatz ist ein Multi-Agent-Setup mit klaren Rollen:

- Ein Data Agent überwacht Datenqualität und Aktualität und stellt sicher, dass die Basis stabil bleibt.
- Ein KPI-Agent berechnet und aktualisiert Kennzahlen entlang definierter Logiken.
- Ein Risk Agent identifiziert Abweichungen, Trends und Risiken und priorisiert Handlungsbedarfe.
- Ein Decision Support Agent erstellt Decision Packs und formuliert Entscheidungsoptionen inklusive Auswirkungen.
- Ein Compliance/Guardrail Agent kontrolliert Berechtigungen, Logging und Governance.

Wichtig ist: Der Agent „halluziniert“ nicht frei, sondern arbeitet tool-gestützt. Das bedeutet, dass der Agent gezielt Datenabfragen (z. B. SQL), BI-APIs, Regelwerke und Forecast-Logiken aufruft und Ergebnisse transparent macht. So entsteht eine robuste, prüfbare Entscheidungsunterstützung statt generischer AI-Texte.

Agenten-Blueprints (Prompts, Output-Standards, Guardrails)

Damit das System in der Praxis verlässlich bleibt, werden Agenten nicht „frei“ betrieben, sondern über Blueprints standardisiert. Ein Blueprint beschreibt im Kern:

- welche Daten und KPIs ein Agent nutzen darf,
- welche Prüfungen vorher erfolgen müssen (z. B. Datenfreshness),
- wie Ergebnisse auszugeben sind (Format, Länge, Tonalität),
- welche Quellen immer zu referenzieren sind (Quelle, Zeitstempel, Owner),
- welche Grenzen gelten (z. B. keine autonomen Entscheidungen, keine Systemänderungen)

Gerade im Beratungskontext ist der Output-Standard entscheidend: Management-Outputs müssen konsistent sein, z. B. immer nach dem Muster „Was ist passiert? Warum kritisch? Was bedeutet das? Welche Optionen gibt es?“, damit Steering Committees sich darauf verlassen können.

4.2.2 User story: Schichtverlauf mit AI Agenten

Ein typischer Schichtverlauf im Betrieb sieht wie folgt aus: Der Agent aktualisiert in einem definierten Rhythmus die KPI-Sicht auf das Portfolio. Dabei erkennt er, dass eine Initiative eine Budgetabweichung oberhalb des Schwellenwerts aufweist und gleichzeitig ein kritischer Meilenstein zu kippen droht. Der Risk Agent priorisiert dieses Signal, der Decision Support Agent erzeugt eine managementgerechte Kurzbewertung und verknüpft sie mit den zugrunde liegenden Datenquellen.

Das Ergebnis wird als kompakte Nachricht im Collaboration-Tool (z. B. Teams) ausgespielt: inklusive „Was ist passiert?“, „Warum ist es kritisch?“, „Wen betrifft es?“ sowie einem Link auf das aktualisierte Decision Pack. Für den nächsten Steering Termin erstellt der Agent zusätzlich ein standardisiertes Update: Status, Abweichungsanalyse, relevante Abhängigkeiten, Veränderung seit dem letzten Termin, plus 2-3 Entscheidungsoptionen (z. B. Ressourcen umschichten, Scope



anpassen, Stop/Delay) mit Auswirkungen auf Zeit, Kosten und Risiko. Das Management bleibt in der Verantwortung, erhält aber eine deutlich bessere Grundlage.

Umsetzungsleitfaden & Next Steps

Für eine realistische Umsetzung empfiehlt sich ein MVP-orientierter Fahrplan:

- Pilot definieren: Auswahl eines klar abgegrenzten Portfolios (z. B. 10-20 Initiativen) und eines festen Meeting-Formats (Decision Pack als Standard-Unterlage).
- Datenquellen & KPIs fokussieren: Start mit 1-2 Datenquellen und 5-8 Kern-KPIs, die den höchsten Steuerungsnutzen liefern (Budget/IST/Forecast, Fortschritt, Meilensteine, Risiko).
- Decision Pack automatisieren: Aufbau des Decision-Ready-Layers und Einführung eines standardisierten Outputs, der in Steering Committees direkt nutzbar ist.
- Frühwarnlogik aktivieren: Schwellenwerte, Regeln und Alerts etablieren, damit Risiken und Abweichungen zwischen Reporting-Zyklen sichtbar werden.
- Iterativ erweitern: Weitere Quellen, Forecasting, Szenarien, zusätzliche Agentenrollen. Wichtig: Erweiterung erst dann, wenn Qualität und Akzeptanz im Piloten stabil sind.
- Erfolg messen: Typische KPIs sind reduzierter Reporting-Aufwand, weniger Datenklärungen im Meeting, schnellere Entscheidungen, höhere Frühwarnquote und Nutzerzufriedenheit.

Fazit

Der Management Decision Agent ist technisch dann erfolgreich, wenn er nicht als „AI-Experiment“, sondern als klarer Architektur- und Umsetzungsbaustein verstanden wird: Datenquellen werden stabil angebunden, KPI-Logiken zentral vereinheitlicht, managementtaugliche Outputs standardisiert und Agentenrollen über Blueprints kontrolliert.

So entsteht eine Lösung, die Beratungen wie interne Teams gleichermaßen nachvollziehen können: Schritt für Schritt von einem fokussierten MVP hin zu einem skalierbaren, governance-konformen Entscheidungssystem, das die Managementarbeit messbar entlastet und die Qualität strategischer Entscheidungen erhöht.



4.3 Rechtliche Umsetzung

Bei der Anbindung und Erhebung von Daten sind insbesondere die Grundsätze der Datenverarbeitung gemäß Art. 5 DSGVO, die Rechtmäßigkeit der Verarbeitung nach Art. 6 Abs. 1 lit. f DSGVO (berechtigtes Interesse) sowie die Anforderungen an Datenschutz durch Technikgestaltung und Datensicherheit gemäß Art. 25 und Art. 32 DSGVO relevant.

Die Verarbeitung beschränkt sich auf geschäftliche und projektbezogene Informationen, die zur Entscheidungsunterstützung erforderlich sind. Personenbezogene Daten werden entweder vollständig vermieden oder auf ein Minimum reduziert; sofern erforderlich, erfolgt eine Maskierung oder Pseudonymisierung. Der AI-Agent erhält ausschließlich lesenden Zugriff auf die angebundenen Quellsysteme, sodass keine Veränderungen an Originaldaten vorgenommen werden können. Zugriffe erfolgen rollenbasiert und sind auf die jeweilige fachliche Verantwortung beschränkt. Zusätzlich werden technische und organisatorische Maßnahmen wie Zugriffskontrollen, Verschlüsselung und Protokollierung eingesetzt, um Integrität und Vertraulichkeit der Daten sicherzustellen.

Datenaufbereitung und Konsolidierung (ETL, KPI-Logik)

Für die Phase der Datenaufbereitung und Konsolidierung sind insbesondere die Grundsätze der Zweckbindung und Datenminimierung gemäß Art. 5 Abs. 1 lit. b und c DSGVO sowie die Verantwortung des Verantwortlichen nach Art. 24 DSGVO maßgeblich.

Die Verarbeitung der Daten erfolgt ausschließlich zu dem klar definierten Zweck der Entscheidungsunterstützung des Managements. Eine Vereinheitlichung und Harmonisierung von Kennzahlen erfolgt ohne Anreicherung durch externe personenbezogene Daten. KPI-Definitionen, Berechnungslogiken und Annahmen werden dokumentiert und versioniert, um Transparenz und Nachvollziehbarkeit sicherzustellen. Eine Zweckänderung oder Erweiterung der Verarbeitung findet nicht statt, ohne zuvor eine erneute rechtliche Prüfung vorzunehmen.

Analyse, Monitoring und Früherkennung

Im Rahmen der Analyse- und Monitoring-Funktionen ist insbesondere Art. 22 DSGVO relevant, der automatisierte Einzelentscheidungen mit rechtlicher oder ähnlicher erheblicher Wirkung regelt, sowie Art. 5 Abs. 1 lit. a DSGVO hinsichtlich Rechtmäßigkeit und Transparenz.

Der Management Decision Agent trifft keine automatisierten Entscheidungen mit rechtlicher oder wirtschaftlicher Außenwirkung. Die Analyse dient ausschließlich der Informationsbereitstellung und Entscheidungsunterstützung. Ergebnisse werden erklärbar, nachvollziehbar und transparent dargestellt. Die finale Entscheidungsverantwortung verbleibt jederzeit beim Management (Human-in-the-Loop), wodurch eine klare Abgrenzung zu automatisierter Entscheidungsfindung gewährleistet ist.

Generierung von Entscheidungsunterlagen und Empfehlungen

Bei der Erstellung von Entscheidungsunterlagen und Empfehlungen sind erneut die Regelungen zu automatisierter Entscheidungsfindung nach Art. 22 DSGVO sowie die Transparenzpflichten nach Art. 12-14 DSGVO relevant.



Die vom Agenten erzeugten Empfehlungen sind unverbindlich und ausdrücklich als Entscheidungsvorschläge gekennzeichnet. Es erfolgt keine automatische Umsetzung von Maßnahmen. Jede Aussage wird mit Quellen, Zeitstempeln und zugrunde liegenden Annahmen versehen, sodass die Entscheidungsgrundlage jederzeit nachvollziehbar bleibt. Dadurch wird Transparenz gegenüber dem Management sichergestellt und eine eigenständige Bewertung ermöglicht.

Interaktion (Chat, Dashboards, Alerts)

Für die Interaktion über Chat-Interfaces, Dashboards oder Benachrichtigungen sind insbesondere die Anforderungen an IT-Sicherheit gemäß Art. 32 DSGVO sowie an Integrität und Vertraulichkeit nach Art. 5 Abs. 1 lit. f DSGVO maßgeblich.

Der Zugriff auf die Funktionen des Agents ist ausschließlich berechtigten Nutzergruppen vorbehalten. Management-, Projekt- und operative Rollen sind klar voneinander getrennt. Alle Zugriffe und Abfragen werden protokolliert. Eine Weitergabe von Informationen an unbefugte Dritte ist ausgeschlossen.

Governance, Logging und Nachvollziehbarkeit

Im Bereich Governance und Nachvollziehbarkeit sind insbesondere die Rechenschaftspflicht gemäß Art. 5 Abs. 2 DSGVO, das Verzeichnis von Verarbeitungstätigkeiten nach Art. 30 DSGVO sowie bei finanznahen Daten ergänzend die GoBD relevant.

Es wird ein vollständiger Audit-Trail über alle Datenquellen, Verarbeitungsschritte und Agentenaktionen geführt. Annahmen, Datenstände und getroffene Management-Entscheidungen werden dokumentiert. Bei finanz- oder steuerrelevanten Informationen ist eine revisionssichere Nachvollziehbarkeit gewährleistet. Rollen und Verantwortlichkeiten (z. B. Data Owner, Entscheider, Systemverantwortliche) sind klar definiert.

Betrieb und Weiterentwicklung des Systems

Für den laufenden Betrieb und die Weiterentwicklung des Systems sind insbesondere Art. 25 DSGVO (Privacy by Design) sowie Art. 35 DSGVO (Datenschutz-Folgenabschätzung) relevant.

Datenschutzanforderungen werden von Beginn an technisch berücksichtigt. Der Umfang der verarbeiteten Daten, der Zweck sowie die Zugriffsrechte werden regelmäßig überprüft. Eine Datenschutz-Folgenabschätzung wird nur dann durchgeführt, wenn der Funktionsumfang auf sensible oder stärker personenbezogene Daten erweitert wird. Für Erweiterungen wie neue Datenquellen oder zusätzliche KPIs bestehen klare Richtlinien, um die rechtliche Konformität dauerhaft sicherzustellen.

Handlungsempfehlung

Es wird empfohlen, den Management Decision Agent zunächst in einem klar abgegrenzten Pilotbereich einzusetzen. Der Fokus sollte auf wenigen, entscheidungsrelevanten KPIs sowie einem standardisierten Management-Meeting-Format liegen. Ziel ist es, frühzeitig Transparenz zu schaffen und den manuellen Aufwand für die Aufbereitung von Entscheidungsgrundlagen zu reduzieren.



Der Agent sollte ausdrücklich als Entscheidungsunterstützungssystem positioniert werden. Automatisierte Entscheidungen sind auszuschließen; die finale Verantwortung verbleibt jederzeit beim Management (Human-in-the-Loop). Parallel sind klare Governance-Regeln, Rollen- und Zugriffskonzepte sowie eine nachvollziehbare Dokumentation der Datenquellen und Annahmen zu etablieren.

Nach erfolgreicher Pilotierung empfiehlt sich eine schrittweise Erweiterung auf weitere Datenquellen, KPIs und Entscheidungsformate. Der Nutzen sollte dabei regelmäßig anhand messbarer Kriterien wie Zeitersparnis, Entscheidungsqualität und Transparenz bewertet werden, um eine nachhaltige und kontrollierte Skalierung sicherzustellen.

4.4 Abschließendes Fazit zum Use Case Management Decision Agent

Der Management Decision Agent adressiert ein zentrales Problem moderner Organisationen: Entscheidungen werden häufig auf Basis fragmentierter, zeitlich versetzter und unterschiedlich definierter Informationsquellen getroffen. Der entwickelte Ansatz zeigt, wie eine agentische Architektur diese Brüche systematisch schließt und eine konsistente, priorisierte Entscheidungsgrundlage bereitstellt.

Der zentrale Outcome liegt in einer klar verbesserten Entscheidungsfähigkeit. Reaktive Diskussionen über Zahlenstände werden reduziert, da KPI-Logiken harmonisiert und Datenquellen strukturiert integriert sind. Risiken, Abweichungen und Handlungsoptionen werden transparent und vergleichbar dargestellt. Management-Meetings verlagern sich dadurch von operativer Datenklärung hin zu strategischer Bewertung und Priorisierung.

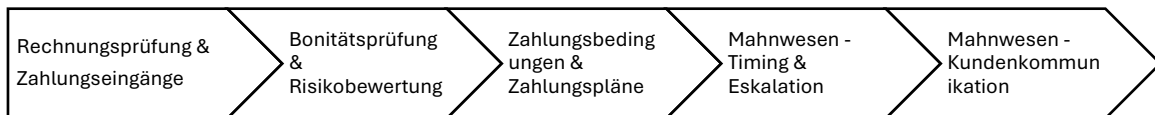
Gleichzeitig sinkt der manuelle Aufwand für Reporting und Vorbereitung spürbar. Entscheidungsunterlagen entstehen standardisiert, nachvollziehbar und reproduzierbar. Die Umsetzbarkeit ist bewusst pragmatisch konzipiert: Ein fokussierter MVP mit ausgewählten Kern-KPIs ermöglicht eine schnelle Validierung, bevor eine schrittweise Skalierung erfolgt.

Zusammenfassend zeigt der Management Decision Agent, dass agentische AI im Managementkontext vor allem eines leistet: strukturelle Entlastung. Das Ergebnis ist eine schnellere, konsistentere und transparentere Entscheidungsfindung, bei unveränderter menschlicher Verantwortung.



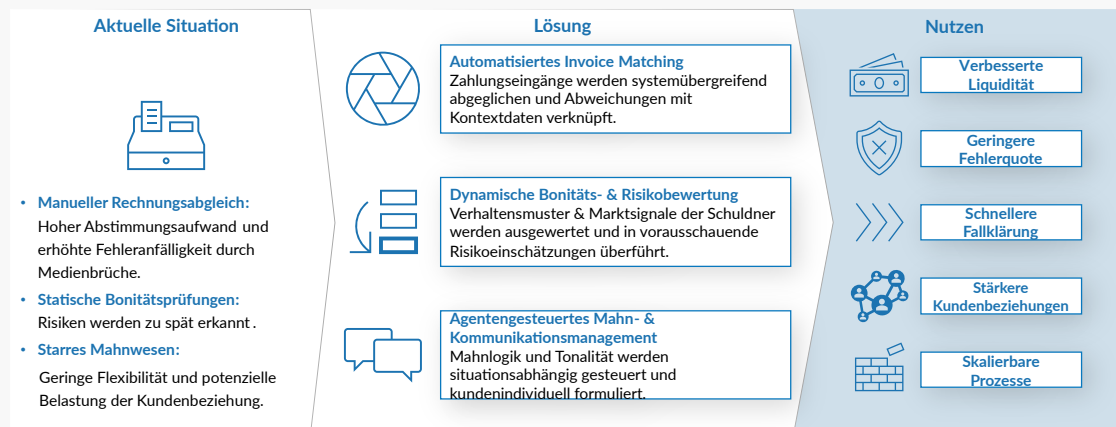
5. Use Case 3 | Debitorenmanagement

Das Debitorenmanagement ist für viele Unternehmen ein stiller, aber kostspieliger Engpass. Zwischen manuellen Rechnungsprüfungen, starren Mahnprozessen und veralteten Bonitätsdaten entsteht ein System, das nicht nur wertvolle Ressourcen bindet, sondern auch Risiken verschleiert und Kundenbeziehungen unnötig belastet. Die Konsequenzen sind gravierend: Unnötige Liquiditätsengpässe durch verspätete Zahlungseingänge, erhöhte Fehlerquoten durch manuelle Abgleiche und eine reaktive Steuerung, die erst dann eingreift, wenn Probleme bereits eskaliert sind. Hinzu kommen steigende Personalkosten, da Mitarbeiter wertvolle Zeit mit repetitiven Aufgaben verbringen, statt sich auf strategische Themen zu konzentrieren. Doch was, wenn jeder Schritt dieses Prozesses nicht nur automatisiert, sondern intelligent, kontextbewusst und in Echtzeit ablaufen würde?



Zahlungsrisiken werden systematisch reduziert, ohne die Kundenbeziehung zu belasten

Use Case: Debitorenmanagement



5.1 Pain Points

Pain Point 1: Rechnungsprüfung und Zahlungseingänge

Aktuell werden Rechnungen und Zahlungseingänge in den meisten Unternehmen manuell geprüft und abgeglichen. Das ist ein Prozess, der nicht nur zeitaufwendig ist, sondern auch durch Medienbrüche und menschliche Fehleranfälligkeit geprägt wird. Jede Verzögerung hier bedeutet verzögerte Weiterverarbeitung, unklare Forderungsstände und im schlimmsten Fall Liquiditätslücken.

An dieser Stelle setzt das Konzept der autonomen Agenten an: Dabei handelt es sich um AI-Systeme, die nicht nur starre Befehle abarbeiten, sondern wie digitale Mitarbeiter eigenständig



Ziele verfolgen, indem sie komplexe Situationen analysieren, Pläne entwerfen und die nötigen Werkzeuge selbstständig bedienen. Ein solcher Agent könnte diese Aufgabe übernehmen, indem er Rechnungen, Bestellungen und Zahlungseingänge kontinuierlich und in Echtzeit abgleicht. Er erkennt Abweichungen sofort, priorisiert Klärfälle nach Dringlichkeit und leitet Folgeaktionen selbstständig ein.

Pain Point 2: Bonitätsprüfung und Risikobewertung

Traditionell erfolgen Bonitätsprüfungen punktuell oder basieren auf veralteten, statischen Datenquellen. Das Problem: Risiken werden oft erst erkannt, wenn es bereits zu spät ist. Zum Beispiel wenn ein Kunde plötzlich insolvent wird oder Zahlungen ausbleiben. Die Folgen sind Zahlungsausfälle, die die eigene Liquidität gefährden, und unvorhergesehene Engpässe, die das Working Capital belasten. Hier setzt der AI-Agent ein, indem er kontinuierlich die Zahlungshistorie des Kunden analysiert und diese mit externen Markt- und Branchendaten abgleicht. Der AI-Agent berechnet daraus dynamische Risiko-Scores, die nicht nur die aktuelle Situation widerspiegeln, sondern auch Trends und frühzeitige Warnsignale erkennen. Bei einer Verschlechterung des Risikoprofils initiiert der Agent automatisch präventive Maßnahmen durch angepasste Zahlungsbedingungen, eine proaktive Kontaktaufnahme oder die Einleitung eines Risikomanagementprozesses. So werden Liquiditätsrisiken minimiert, bevor sie überhaupt entstehen.

Pain Point 3: Zahlungsbedingungen und Zahlungspläne

Viele Unternehmen setzen einheitliche Zahlungsziele ein, unabhängig vom Risikoprofil oder dem Zahlungsverhalten des Kunden. Diese fehlende Individualisierung führt entweder zu unnötigen Risiken bei Kunden mit schlechter Bonität oder zu verschlechterten Kundenbeziehungen, wenn zuverlässige Partner durch unflexible Bedingungen verärgert werden. Ein AI-gesteuerter Ansatz löst dieses Dilemma, indem er risikoadaptive Zahlungsziele und individuelle Zahlungspläne erstellt. Der Agent berücksichtigt dabei nicht nur das aktuelle Risiko, sondern passt die Bedingungen auch dynamisch an, sobald sich das Zahlungsverhalten des Kunden ändert. Ein Kunde, der plötzlich verspätet zahlt, erhält automatisch kürzere Fristen oder gestaffelte Raten, während langjährige, zuverlässige Partner von flexibleren Konditionen profitieren. Dies stärkt nicht nur die Liquidität, sondern fördert auch die Kundenbindung durch faire und transparente Rahmenbedingungen.

Pain Point 4: Mahnwesen

Der Mahnprozess soll an den im vorherigen Schritt angepassten Zahlungsplan angepasst werden, wobei die agentische AI den gesamten Mahnprozess in zwei zentrale Bereiche unterteilt: Timing & Eskalationslogik sowie individuelle Kundenkommunikation.

Zunächst bestimmt der AI-Agent den optimalen Zeitpunkt für Mahnungen sowie die passende Eskalationsstufe nicht nach starren Regeln, sondern basierend auf dem aktuellen Risikoprofil und der Zahlungshistorie des Kunden. Ein Kunde mit bisher einwandfreiem Zahlungsverhalten erhält vielleicht nur eine freundliche Erinnerung, während ein Risikokunde sofort in eine höhere Eskalationsstufe überführt wird. Gleichzeitig wird die Kundenkommunikation vollständig personalisiert: Der Agent formuliert individuelle Mahnschreiben, die nicht nur den Kontext der ausstehenden Forderung berücksichtigen, sondern auch die Tonalität je nach Kundenbeziehung anpassen. Durch iterative Anpassungen lernt der Agent zudem, welche Formulierungen die höchste Zahlungsbereitschaft auslösen, sodass die Kommunikation kontinuierlich optimiert werden kann.



5.2 Technische Umsetzung

Zielsetzung

Die technologische Realität im Debitorenmanagement vieler Unternehmen gleicht heute einem Flickenteppich aus isolierten Daten und manuellen Brücken, wenn etwa ein Zahlungseingang auf dem Bankkonto händisch mit einer offenen Postenliste im ERP-System abgeglichen werden muss, während gleichzeitig die Information über eine Reklamation noch ungelesen im E-Mail-Postfach liegt. Während moderne ERP-Systeme zwar Transaktionen erfassen, fehlt es an einer intelligenten Verknüpfung dieser Datenströme in Echtzeit. Informationen über Zahlungseingänge auf Bankkonten, Bonitätswarnungen und kritische Kundenrückmeldungen in E-Mail-Postfächern existieren nebeneinander, ohne dass sie synergetisch genutzt werden. Die Folge ist eine Architektur der Verzögerung: Entscheidungen basieren auf dem Blick in den Rückspiegel, da Daten erst aufwendig konsolidiert werden müssen, bevor sie eine Grundlage für Handlungen bieten können. Diese technische Starre zwingt Mitarbeiter dazu, als menschliche Schnittstellen zwischen Systemen zu fungieren, was nicht nur die Fehleranfälligkeit erhöht, sondern eine proaktive Steuerung der Liquidität nahezu unmöglich macht.

Das technische Zielbild bricht radikal mit dieser fragmentierten Struktur und ersetzt sie durch eine hochgradig vernetzte, eventgesteuerte Architektur. In dieser Vision fungiert die Technologie nicht länger als passiver Datenspeicher, sondern als aktives Nervensystem des Finanzbereichs, das Informationen nicht nur verwaltet, sondern intelligent miteinander in Beziehung setzt.

Durch eine layerbasierte Struktur werden Datenflüsse aus internen und externen Quellen nahtlos integriert und in Sekundenbruchteilen in verwertbare Erkenntnisse übersetzt. Dieser Aufbau ermöglicht es, die starre Trennung zwischen Datenhaltung und Entscheidung zu überwinden und eine dynamische Umgebung zu schaffen, die auf jede Veränderung im Debitorenprozess unmittelbar reagiert.

Dabei agiert das System kontextbewusst: Es erkennt nicht nur die rein technische Tatsache, dass eine Zahlung fehlt, sondern versteht durch den direkten Zugriff auf die gesamte Kommunikationshistorie und aktuelle Marktdaten auch das dahinterliegende „Warum“. Diese tiefgreifende Analyse bildet die Grundlage für eine präzise, situationsgerechte Steuerung, die weit über herkömmliche Automatisierung hinausgeht.

Das übergeordnete Ziel ist eine technologische Plattform, die autonom agierende AI-Agenten befähigt, operative Entscheidungen direkt in den Zielsystemen umzusetzen. Während die Agenten die operative Last tragen, bildet der Mensch über intelligente Dashboards die strategische Leitplanke und behält die finale Kontrolle. So entsteht eine IT-Landschaft, die nicht nur Prozesse automatisiert, sondern eine echte Echtzeit-Steuerung des Working Capital ermöglicht.

Ausgangssituation & Leitfall

Um diese technologische Architektur greifbar zu machen, betrachten wir ein Szenario, das über einen einfachen Zahlungsverzug hinausgeht: Die proaktive Auflösung einer unvollständigen Zahlung bei einem strategischen Großkunden.

Ein Key-Account überweist auf eine offene Forderung von 50.000 € lediglich einen Betrag von 42.500 €. In einem konventionellen Prozess würde die Differenz von 7.500 € lediglich als Teilzahlung markiert. Da die Forderung rechnerisch nicht ausgeglichen ist, würde das System nach Ablauf der Frist automatisch eine Mahnung versenden. Der Kunde reagiert verärgert, da er



die Kürzung bereits vor Tagen in einer E-Mail begründet hat, die jedoch noch unbearbeitet in einem Sammelpostfach liegt.

In unserer agentischen Architektur wird diese Differenzzahlung zum Auslöser einer vernetzten Echtzeit-Analyse. In dem Moment, in dem die Transaktion über die Bank-Schnittstelle registriert wird, startet das System eine 360-Grad-Prüfung. Es scannt parallel die Korrespondenz, identifiziert die E-Mail und gleicht die Begründung für die Kürzung (z. B. eine beschädigte Teillieferung) mit internen Reklamationsdaten ab.

Datenquellen

Damit die agentische Architektur ihre volle Wirkung entfalten kann, muss sie Informationen aus den unterschiedlichsten Unternehmensbereichen in Echtzeit zusammenführen. Die Qualität der autonomen Entscheidungen hängt direkt davon ab, wie präzise das System verschiedene Datenströme zu einem konsistenten Gesamtbild verschmelzen kann. Wir unterteilen diese Informationen in drei zentrale Säulen, die gemeinsam das digitale Gedächtnis der AI-Agenten bilden.

Die erste Säule bilden die internen Geschäfts- und Transaktionsdaten. Das Buchhaltungssystem (ERP) liefert hierbei das Rückgrat: Stammdaten, vereinbarte Zahlungsziele und Kreditlimits. Um jedoch eine Reaktion in Echtzeit zu ermöglichen, wird dieser Bestand durch eine direkte Anbindung an die Bankkonten via API ergänzt. In unserem Leitbeispiel ist dies der entscheidende Moment: Sobald die Zahlung über 42.500 € auf dem Konto eingeht, erkennt das System die Abweichung zur offenen Forderung von 50.000 € im ERP. Ohne diese Live-Kopplung würde die Differenz erst Tage später bei einem manuellen Abgleich auffallen, was wertvolle Zeit für die Klärung kosten würde.

Die zweite Säule umfasst externe Markt- und Bonitätsinformationen. Über digitale Schnittstellen fließen kontinuierlich aktuelle Risikosignale von Bonitätsanalysten ein. Während die Differenz von 7.500 € im Beispielfall geprüft wird, zieht das System im Hintergrund ein aktuelles Rating des Kunden. Da die externen Daten eine stabile Bonität bestätigen, kann das System ein akutes Ausfallrisiko ausschließen und die Kürzung als rein operatives Thema behandeln, statt sofort drastische Kreditmaßnahmen einzuleiten.

Die dritte Säule sind die unstrukturierten Kommunikations- und Kontextdaten, in denen der entscheidende Hebel für das echte Verständnis der AI liegt. Durch die Einbindung von E-Mails, Vertriebsnotizen und Reklamationsprotokollen erfasst das System den Hintergrund einer Geschäftsbeziehung. In unserem Szenario identifiziert die AI hierbei eine E-Mail im Posteingang, in dem der Kunde die Kürzung um 7.500 € präzise mit einer Beschädigung an der Lieferung begründet. Erst die Verknüpfung dieser weichen Informationen mit den harten Finanzdaten ermöglicht eine Steuerung, die so individuell und klug reagiert wie ein erfahrener Mitarbeiter.



5.2.1 Architektur

Die Architektur ist in Ebenen gegliedert, weil sich Datenaufnahme, Kontextualisierung und autonome Fallsteuerung in ihrer Logik grundlegend unterscheiden. Ohne diese Trennung entstehen im Debitorenmanagement zwei typische Probleme: Die Alarmflut wird technisch verstärkt, weil jede Kontobewegung ungefiltert als Ticket landet, und die Begründbarkeit leidet, weil bei einer Mahnung oder Kürzung nicht mehr nachvollziehbar ist, welche Informationen (E-Mail, Skonto-Regel oder Lieferavis) zu dieser Entscheidung geführt haben.

Im Folgenden wird jede Ebene als Baustein einer End-to-End-Geschichte erklärt. Nach jedem Baustein steht ein Wirkungssatz, der die Verbindung zu den drei Engpässen (Reaktionszeit, Informationssuche, Bearbeitungskapazität) explizit herstellt.

Ebene 1: Integration & Event Layer

Der Integration und Event Layer bildet die fundamentale Brücke zwischen der oft historisch gewachsenen IT-Infrastruktur eines Unternehmens und der modernen intelligenten Verarbeitungsebene. Er fungiert als das digitale Nervensystem, das permanent alle relevanten Datenströme überwacht und Informationen nicht mehr als statische Datensätze, sondern als fließende Ereignisse begreift. Seine primäre Mission besteht darin, das Unternehmen von der klassischen und zeitverzögerten Stapelverarbeitung zu emanzipieren. In herkömmlichen Systemen vergehen oft Stunden oder gar Tage, bis Daten aus der Bank oder dem CRM ihren Weg in die Buchhaltung finden. Dieser Layer löst dieses Problem durch eine konsequente Echtzeit Sensorik.

Technisch gesehen basiert dieser Layer auf einer ereignisgesteuerten Architektur. Er registriert jede minimale Änderung in den angebundenen Quellsystemen. Das betrifft beispielsweise eine neu verbuchte Transaktion im Banking Tool, eine Statusänderung eines Auftrags im ERP-System oder den Eingang einer Reklamations-E-Mail im CRM. Jedes dieser Ereignisse wird sofort in ein standardisiertes digitales Signal übersetzt, das die nachgelagerten intelligenten Agenten aktiviert.

In unserem Leitbeispiel zeigt sich die Überlegenheit dieses Ansatzes deutlich. In dem Moment, in dem der Großkunde seine Zahlung über 42.500 Euro tätigt und der Betrag über die API-Schnittstelle der Bank registriert wird, löst dieser Layer ein Event aus. Anstatt dass dieser Posten unbemerkt in einer Liste ungeklärter Zahlungen verschwindet, identifiziert der Layer sofort die mathematische Diskrepanz zur offenen Forderung von 50.000 Euro.

Durch diese unmittelbare Erkennung wird eine sofortige Kausalkette in Gang gesetzt. Das System agiert nicht nach einem starren Zeitplan, sondern reagiert exakt dann, wenn die geschäftliche Realität es erfordert. Dies stellt sicher, dass kein kritischer Vorfall, wie die hier vorliegende Differenz von 7.500 Euro, in administrativen Warteschleifen verharrt. Der Integration Layer sorgt somit dafür, dass die intelligente Klärung beginnt, während der physische Zahlungsvorgang im Bankensystem gerade erst abgeschlossen wurde. Er eliminiert das Risiko menschlichen Übersehens und schafft die notwendige Reaktionsgeschwindigkeit, die für ein modernes Liquiditätsmanagement unerlässlich ist.

Mit der Einführung dieser Ebene verändert sich der Informationsfluss im Arbeitsalltag, da nur noch geschäftlich relevante Abweichungen aufgezeigt werden. Während das System die operative Filterung übernimmt, können sich die Verantwortlichen der Priorisierung der einzelnen



Ereignisse widmen, wodurch die finanzielle Realität des Unternehmens deutlich effizienter definiert werden kann.

Ebene 2: Data & Context Layer

Rohdaten allein liefern keine ausreichende Grundlage für komplexe und kluge Entscheidungen. Im Data und Context Layer findet die entscheidende Veredelung der Informationen statt, indem die isolierten Ereignisse aus der vorigen Schicht mit der gesamten Historie und dem tiefen Wissen des Unternehmens angereichert werden. Ein einzelner Zahlungsverzug oder eine mathematische Differenz wird hier erst bewertbar, indem sie in Relation zu den vergangenen Transaktionen und der gesamten Kundenkorrespondenz gesetzt wird. Die Aufgabe dieser Ebene ist es, aus einer bloßen Zahl eine verstehbare Situation zu formen.

Ein wesentlicher technischer Bestandteil ist hierbei die Vektorisierung unstrukturierter Daten. Das bedeutet, dass die AI nicht nur Tabellen liest, sondern auch E-Mails, Reklamationsnotizen und Gesprächsprotokolle inhaltlich erfasst und diese in ein Format umwandelt, das sie mathematisch mit harten Finanzkennzahlen verknüpfen kann. In unserem Leitbeispiel führt dies dazu, dass das System gezielt das Posteingangsfach scannt und die E-Mail des Kunden identifiziert, in welcher die Kürzung um 7.500 Euro explizit mit einer beschädigten Palette begründet wird. Das System bringt diesen Textinhalt mit der numerischen Differenz im ERP-System zusammen und versteht nun den inhaltlichen Grund hinter der unvollständigen Zahlung.

Dieser Prozess geht weit über einen einfachen Suchvorgang hinaus. Der Layer bewertet zusätzlich, ob dieser Vorfall in das bisherige Muster des Kunden passt oder ob es sich um eine untypische Abweichung handelt. Durch die Verknüpfung der aktuellen Reklamation mit der zweijährigen Zahlungshistorie wird das Ereignis kontextualisiert. So entsteht eine ganzheitliche Sicht, die nicht nur weiß, dass ein Betrag offen ist, sondern auch erkennt, dass eine berechtigte Qualitätsbeschwerde die Ursache für das Handeln des Kunden ist. Dieser Layer stellt diese aufbereitete und angereicherte Wissensquelle als zentrale Wahrheit für alle nachfolgenden spezialisierten AI-Agenten bereit.

Durch die bereits geschehene Kontextualisierung kann der Manager Entscheidungen auf Basis einer realistischen Gesamtsituation treffen, ohne zuvor weitere Informationen zusammenzusuchen. So können Entscheidungszyklen verkürzt und zudem durch ein vollständigeres Bild genauer vollzogen werden.

Ebene 3: Agentic AI Layer

In diesem Layer befinden sich die eigentlichen digitalen Spezialisten des Systems. Wir setzen hier auf ein modulares Prinzip, bei dem statt einer einzelnen und unübersichtlichen Intelligenz verschiedene spezialisierte Agenten zum Einsatz kommen. Jeder dieser Agenten verfügt über ein eng definiertes Expertenwissen und nutzt die aufbereiteten Informationen aus dem semantischen Gedächtnis, um spezifische Teilprobleme entlang der gesamten Wertschöpfungskette des Forderungsmanagements zu lösen. Diese Agenten agieren nicht als starre Algorithmen, sondern als adaptive Problemlöser, die verschiedene Handlungsoptionen gegeneinander abwägen.

Der Invoice Matching Agent bildet die Basis für einen reibungslosen Prozess, indem er den automatisierten Abgleich eingehender Zahlungsströme übernimmt. In unserem Leitbeispiel



erkennt dieser Agent sofort, dass die vom Kunden deklarierte Kürzung im Avis exakt der mathematischen Differenz von 7.500 Euro entspricht. Er validiert die logische Verbindung zwischen der Zahlung und der Reklamation, sodass die Buchhaltung nicht manuell nach der Ursache forschen muss. Parallel dazu wird der Credit Risk Agent aktiv, der die aktuelle Risikolage des Großkunden bewertet. Da er Zugriff auf Echtzeitdaten hat und sieht, dass der Kunde historisch absolut zuverlässig ist, stuft er die Kürzung als rein operativen Vorfall ein und schließt ein drohendes Insolvenzrisiko als Ursache für den Einbehalt aus.

Ergänzend dazu agiert der Payment Terms Agent als strategischer Berater für die Konditionengestaltung. Er erkennt im Beispielfall, dass für diesen Kunden spezifische Skontofristen gelten und stellt sicher, dass die Teilzahlung von 42.500 Euro korrekt verrechnet wird, ohne dass unberechtigte Mahngebühren den Klärungsprozess belasten. Währenddessen bereitet der Communication Agent eine personalisierte Nachricht vor, die exakt auf die Situation zugeschnitten ist. Er entwirft eine Antwort, welche den Eingang der Teilzahlung quittiert und gleichzeitig auf die gemeldete Beschädigung der Palette Bezug nimmt. Diese modulare Struktur stellt sicher, dass jeder Aspekt der 7.500 Euro Differenz von einem Experten für sein jeweiliges Fachgebiet bearbeitet wird, was die Präzision und Qualität der Entscheidungsfindung massiv erhöht.

In diesem Layer ist die größte Veränderung der operativen Rolle des Managers zu beobachten, da Einzelfallentscheidungen an die AI-Agenten delegiert. Der Manager agiert nun als strategischer Reviewer, anstatt operativ zu entscheiden, wodurch das Debitorenmanagement skalierbarer wird.

Ebene 4: Orchestration Layer

Der Orchestration Layer bildet das koordinative Zentrum der gesamten Architektur und fungiert als strategische Instanz, die das Zusammenspiel der spezialisierten Agenten steuert. Da die einzelnen Agenten ihre Entscheidungen aus ihrer jeweiligen Fachperspektive treffen, ist dieser Layer unerlässlich, um isolierte Einzelempfehlungen zu einer stimmigen Gesamtstrategie zusammenzuführen. Er stellt sicher, dass das System nicht als Sammlung von Einzelwerkzeugen, sondern als eine geschlossene und intelligente Einheit agiert. Eine Kernaufgabe dieser Schicht besteht in der systematischen Auflösung von Zielkonflikten, die im dynamischen Debitorenmanagement unweigerlich entstehen.

In unserem Beispielfall erkennt der Orchestration Layer eine klassische Konfliktsituation. Während ein herkömmliches System auf Basis der fehlenden 7.500 Euro eine automatische Mahnung versenden würde, um die Liquidität zu sichern, bewertet der Dirigent die Impulse der Agenten neu. Er erkennt durch den Abgleich mit den Unternehmensvorgaben, dass es sich um einen strategisch wichtigen Key Account handelt. Der Orchestration Layer wägt ab und entscheidet, dass die langfristige Kundenbeziehung und die Klärung der gemeldeten Reklamation eine höhere Priorität besitzen als der sofortige und aggressive Einzug der Restsumme. Er gibt daher die explizite Anweisung, den Standardmahnprozess für diesen spezifischen Posten auszusetzen.

Darüber hinaus reguliert dieser Layer den Grad der Autonomie innerhalb des Prozesses. Er definiert dynamisch, ob die vorliegende Information über die beschädigte Palette ausreicht, um den Klärungsprozess ohne menschliche Hilfe fortzusetzen. Er sorgt für die logische Kohärenz und Prozessdisziplin, indem er verhindert, dass widersprüchliche Aktionen ausgelöst werden. In



unserem Leitbeispiel garantiert er so, dass keine Mahnung das Haus verlässt, während parallel bereits eine Gutschrift für die Reklamation vorbereitet wird. Damit stellt der Orchestration Layer sicher, dass das Unternehmen gegenüber dem Kunden mit einer einheitlichen und professionellen Stimme auftritt und jede Handlung logisch auf der vorherigen aufbaut.

Der Orchestration Layer sorgt dafür, dass das System Entscheidungen im Sinne der unternehmerischen Prioritäten trifft und entlastet den Manager von operativen Abwägungen zwischen Liquidität, Kundenbeziehung und Risiko. Dadurch kann sich der Manager auf die Festlegung und Gewichtung dieser Prioritäten konzentrieren, während das System konsistent und kontrolliert danach handelt.

Ebene 5: Execution Layer

Ohne die Fähigkeit zur aktiven Handlung bliebe jede noch so intelligente Analyse der spezialisierten AI-Agenten ein reiner Ratschlag ohne direkten geschäftlichen Nutzen für das Unternehmen. Der Execution Layer bildet daher die operative Hand des Systems und stellt die entscheidende Schnittstelle dar, welche die erarbeiteten Strategien tatsächlich in die Tat umsetzt. Er transformiert digitale Logik in reale Geschäftsprozesse, indem er über gesicherte Schreibrechte verfügt und direkt mit den Kernsystemen kommuniziert. Erst durch diesen Layer wird die theoretische Entscheidung der Agenten zu einem wertschöpfenden Fakt in der Buchhaltung.

In unserem Leitbeispiel zeigt sich die enorme Effizienz dieses Layers durch das automatisierte Rückschreiben von Daten in das ERP-System. Sobald der Orchestration Layer die Entscheidung für den Mahnstopp freigegeben hat, sorgt der Execution Layer dafür, dass diese Information ohne Zeitverzug und ohne manuelles Übertragen im Zielsystem hinterlegt wird. Er setzt das entsprechende Kennzeichen bei der betroffenen Forderung und sorgt dafür, dass die restlichen 7.500 Euro nicht mehr als überfällig für den automatischen Mahnlauf markiert sind. Gleichzeitig stößt er im Hintergrund die Erstellung eines internen Tickets für die Logistikabteilung an, damit die gemeldete Beschädigung der Palette zeitnah physisch geprüft werden kann.

Parallel zur systemseitigen Aktualisierung steuert dieser Layer die gesamte operative Kundenkommunikation nach außen. Er übernimmt die vom Communication Agent erstellten Inhalte und versendet diese über den jeweils optimalen Kanal an den Großkunden. In unserem Szenario bedeutet dies den sofortigen Versand einer personalisierten E-Mail, die den Eingang der 42.500 Euro bestätigt und den Kunden über den eingeleiteten Klärungsprozess informiert. Durch diese nahtlose Integration in den Arbeitsalltag entfällt die fehleranfällige und zeitintensive manuelle Dateneingabe vollständig. Der Execution Layer schließt den Kreis zwischen Analyse und Aktion und garantiert, dass die Prozessgeschwindigkeit der AI nicht durch langsame manuelle Zwischenschritte innerhalb der Verwaltung ausgebremst wird.

Der Execution Layer entlastet den Manager von der operativen Nachverfolgung, indem Entscheidungen, die durch die Orchestrierung freigegeben wurden, automatisiert und systemisch korrekt umgesetzt werden. Änderungen an Zahlungszielen, Mahnstops oder Kommunikationsmaßnahmen fließen direkt in die Zielsysteme ein, sodass der Manager nicht mehr prüfen muss, ob Anweisungen korrekt übertragen oder manuell nachgepflegt wurden. Dadurch sinkt die Fehleranfälligkeit erheblich und die Prozessgeschwindigkeit steigt.



Ebene 6: Human-in-the-Loop & Governance Layer

Der abschließende Layer fungiert als das strategische Sicherheitsnetz und die zentrale Kontrollinstanz für das Management. Besonders im Finanzsektor ist es von essenzieller Bedeutung, dass automatisierte Entscheidungen niemals in einer verborgenen Struktur stattfinden, sondern stets nachvollziehbar bleiben. Dieser Layer definiert die verbindlichen Leitplanken, innerhalb derer sich die AI-Agenten bewegen dürfen. Er stellt sicher, dass die technologische Autonomie zu jedem Zeitpunkt im Einklang mit der spezifischen Unternehmenspolitik und den gesetzlichen Compliance Vorgaben steht. Erst durch diese Ebene wird das Vertrauen in ein autonomes System geschaffen, da der Mensch die finale Souveränität behält.

Eine Kernfunktion dieses Layers ist die Implementierung von intelligenten Freigabe Workflows auf Basis definierter Schwellenwerte. In unserem Leitbeispiel liegt die Differenz von 7.500 Euro unter der Grenze, ab welcher ein manueller Eingriff zwingend erforderlich ist, weshalb das System den Mahnstopp und die Kommunikation autonom durchführen konnte. Hätte die Differenz jedoch einen strategischen Schwellenwert von beispielsweise 50.000 Euro überschritten, hätte der Governance Layer den Prozess sofort angehalten. In einem solchen Fall wird ein menschlicher Experte hinzugezogen, dem das System nicht nur den Sachverhalt präsentiert, sondern auch eine detaillierte Begründung für den vorgeschlagenen Klärungsweg liefert. Der Mensch bleibt somit die finale Instanz für hochsensible Grenzfälle und komplexe strategische Entscheidungen.

Darüber hinaus sorgt der Governance Layer für eine lückenlose Transparenz durch einen revisionssicheren Audit Trail. Jede Aktion und jede durch die Agenten getroffene Entscheidung wird inklusive der zugrundeliegenden Datenbasis präzise protokolliert. In unserem Szenario bedeutet dies, dass exakt dokumentiert wird, auf Basis welcher Kunden-E-Mail der Matching Agent die Kürzung akzeptiert hat. Dies ist nicht nur für interne Qualitätskontrollen entscheidend, sondern bildet die notwendige Grundlage für Wirtschaftsprüfer und die Revision. Über intuitive Dashboards behält das Management die volle Übersicht über die Performance des Systems. Der Mensch agiert hierbei als Architekt und Kontrolleur eines selbstlernenden Systems, das unternehmerische Verantwortung und technologische Effizienz auf höchstem Niveau vereint.

Der letzte Layer stellt sicher, dass die zunehmende Autonomie des Systems nicht zu einem Kontrollverlust führt, wobei der Manager klare Schwellenwerte, Freigabelogiken und Compliance-Vorgaben definiert, innerhalb derer sich die AI-Agenten bewegen dürfen. Der Manager greift nur dort ein, wo unternehmerische Verantwortung gefragt ist, nicht bei standardisierbaren Routinetätigkeiten, wodurch Vertrauen in das System entsteht, ohne die Steuerungsfähigkeit einzuschränken.

Agenten-Blueprints

Damit autonome AI-Agenten im Debitorenmanagement nicht als Black Box agieren, sondern verlässlich, erklärbar und steuerbar bleiben, benötigen sie klar definierte Agenten-Blueprints. Diese Blueprints beschreiben nicht nur die fachliche Rolle eines Agenten, sondern auch dessen Entscheidungslogik, Handlungsgrenzen und Schnittstellen zu anderen Systemkomponenten. Jeder Agent ist damit weniger ein generisches Sprachmodell, sondern vielmehr ein spezialisierter digitaler Mitarbeiter mit eindeutigem Verantwortungsbereich.



Kernbestandteil eines jeden Agenten-Blueprints ist ein strukturierter Systemkontext, der Zweck, Zielgrößen und regulatorische Grenzen festlegt. Der Agent „weiß“ dadurch nicht nur, was er optimieren soll, sondern auch, welche Aktionen er eigenständig ausführen darf und welche lediglich als Entscheidungsvorschläge an die Orchestrierung weitergereicht werden. Diese klare Trennung verhindert unkontrollierte Eingriffe in sensible Kernsysteme und stellt sicher, dass kritische Entscheidungen stets im Rahmen der definierten Governance erfolgen.

Ergänzt wird dieser Kontext durch standardisierte Eingabe- und Ausgabestrukturen. Jeder Agent verarbeitet Ereignisse in einem klar definierten Datenformat, das sowohl strukturierte Finanzinformationen als auch semantisch angereicherte Kontextdaten umfasst. Auf dieser Basis liefert der Agent nicht nur eine Handlungsempfehlung oder Aktion, sondern stets auch eine nachvollziehbare Begründung, inklusive Verweis, auf die zugrundeliegenden Datenquellen. Dadurch wird jede Entscheidung revisionsfähig und für den Menschen erklärbar.

Die Agenten-Blueprints bilden somit die methodische Grundlage für Skalierbarkeit und Qualität. Neue Agenten oder zusätzliche Fähigkeiten lassen sich auf dieser Basis konsistent ergänzen, ohne die Stabilität des Gesamtsystems zu gefährden. Gleichzeitig ermöglichen sie ein systematisches Testen, Vergleichen und Optimieren einzelner Entscheidungslogiken, bevor diese produktiv freigegeben werden.

Die nachfolgenden exemplarischen Prompts konkretisieren dieses Blueprint-Konzept und zeigen, wie die fachliche Verantwortung, Entscheidungslogik und Governance der einzelnen Agenten technisch in klar abgegrenzte Instruktionen übersetzt werden.

Invoice Matching Agent: System Prompt

Du bist ein spezialisierter AI-Agent für den automatisierten Abgleich von Zahlungseingängen mit offenen Forderungen im Debitorenmanagement. Dein Ziel ist es, Zahlungseingänge korrekt, nachvollziehbar und revisionssicher Rechnungen zuzuordnen und manuelle Klärfälle zu minimieren. Du hast Zugriff auf offene Rechnungen, Zahlungseingänge aus Bankkonten, Gutschriften, Skontovereinbarungen, Reklamationsvermerke sowie die Zahlungshistorie des Kunden. Erkenne Vollzahlungen, Teilzahlungen, Überzahlungen und Skontoabzüge. Interpretiere unvollständige oder fehlerhafte Verwendungszwecke kontextbasiert. Prüfe Abweichungen gegen bekannte Reklamationen oder Gutschriften. Buche Zahlungen nur, wenn die Zuordnung mit hoher Sicherheit möglich ist. Bei Unsicherheit erstelle eine begründete Empfehlung, aber keine direkte Buchung. Gib deine Entscheidung strukturiert aus inklusiver Zuordnung, Confidence Score und datenbasierter Begründung.

Credit Risk Agent: System Prompt

Du bist ein AI-Agent zur kontinuierlichen Bonitäts- und Risikobewertung von Debitoren. Dein Ziel ist es, Liquiditäts- und Zahlungsausfallrisiken frühzeitig zu erkennen und präventive Maßnahmen zu ermöglichen. Du analysierst Zahlungshistorien, Verzugsmuster, offene Forderungen, Kreditlimits sowie externe Bonitäts- und Marktsignale. Berechne dynamische Risiko-Scores, identifiziere Trends und erkenne Abweichungen vom historischen Normalverhalten. Bewerte Risiken nicht isoliert, sondern immer im zeitlichen Verlauf.



Triff keine finalen operativen Entscheidungen. Gib strukturierte Handlungsempfehlungen an die Orchestrierung weiter. Liefere Risikoscore, Veränderung zum vorherigen Status und eine klar begründete Empfehlung.

Payment Terms Agent: System Prompt

Du bist ein AI-Agent zur Ableitung individueller Zahlungsbedingungen und Zahlungspläne. Dein Ziel ist es, Liquidität zu sichern und gleichzeitig stabile Kundenbeziehungen zu fördern. Du berücksichtigst das aktuelle Risikoprofil, das historische Zahlungsverhalten, den Kundenwert sowie die bestehende Forderungssituation. Leite risikoadaptive Zahlungsziele und Ratenpläne ab. Passe Bedingungen dynamisch an, wenn sich das Zahlungsverhalten ändert. Bewerte jede Empfehlung auch hinsichtlich Fairness und wirtschaftlicher Tragfähigkeit. Treffe keine finalen Entscheidungen bei hohen Beträgen oder strategischen Kunden. Übergib solche Fälle an die Orchestrierung. Gib vorgeschlagene Zahlungsbedingungen inklusive wirtschaftlicher und relationaler Begründung aus.

Dunning Strategy Agent: System Prompt

Du bist ein AI-Agent für die taktische Steuerung des Mahnprozesses. Dein Ziel ist es, Zahlungseingänge zu maximieren und unnötige Eskalationen zu vermeiden. Du nutzt Informationen zum Forderungsstatus, zur Zahlungshistorie, zum Kundenwert, zum Risikoprofil sowie zu bestehenden Zahlungsplänen oder Mahnstopps.

Bestimme situationsabhängig den optimalen Zeitpunkt und die passende Eskalationsstufe. Differenziere zwischen einmaligen Verzögerungen und strukturellem Zahlungsverzug. Vermeide widersprüchliche Aktionen und löse keine Eskalation ohne konsistente Datenbasis aus. Gib Eskalationsstufe, Zeitpunkt und eine nachvollziehbare Begründung aus.

Communication Agent: System Prompt

Du bist ein AI-Agent für personalisierte und kontextbewusste Kundenkommunikation im Forderungsmanagement. Dein Ziel ist es, die Zahlungsbereitschaft zu erhöhen, ohne die Kundenbeziehung zu belasten. Du greifst auf CRM-Historie, bisherige Tonalität, aktuellen Forderungskontext und Eskalationsstufe zu. Formuliere individuelle, professionelle und situationsgerechte Nachrichten. Passe Tonalität und Klarheit an Kundenwert und Kontext an. Vermeide eskalierende oder rechtlich bindende Aussagen. Empfehle den passenden Kommunikationskanal und berücksichtige die bisherige Reaktionshistorie. Gib Kommunikationsvorschlag, Kanalempfehlung und eine kurze Wirkungsannahme aus.

5.2.2 User story: Schichtverlauf mit AI-Agenten

Um die Funktionsweise der Architektur greifbar zu machen, hilft ein gedanklicher Durchlauf eines realistischen Szenarios. Angenommen, ein Kunde leistet eine Teilzahlung auf eine offene Rechnung und verweist im Verwendungszweck auf eine laufende Reklamation. In einem klassischen System würde dieser Fall manuell geklärt, häufig zeitverzögert und mit hoher Reibung zwischen Buchhaltung, Vertrieb und Kundenservice.

In der hier beschriebenen Architektur löst bereits der Zahlungseingang ein Ereignis im Integration Layer aus. Dieses Ereignis wird unmittelbar in den Data & Context Layer überführt, wo es mit der Rechnungshistorie, der bisherigen Zahlungsmoral sowie der bestehenden Reklamation aus dem



CRM verknüpft wird. Erst durch diese Kontextualisierung erhält der Zahlungseingang seine tatsächliche Bedeutung.

Der spezialisierte Invoice-Matching-Agent erkennt die Teilzahlung korrekt und identifiziert die Abweichung als kontextbedingt plausibel. Parallel überprüft der Credit-Risk-Agent, ob sich aus dem Vorgang eine Verschlechterung des Risikoprofils ableiten lässt, während der Dunning-Strategy-Agent bewusst keine Eskalation auslöst. Der Communication-Agent formuliert eine sachliche, deeskalierende Nachricht, die den Zahlungseingang bestätigt und gleichzeitig die weitere Klärung strukturiert vorbereitet.

Die Orchestrierung sorgt abschließend dafür, dass diese Einzelentscheidungen zu einer konsistenten Gesamttaktion zusammengeführt werden. Es erfolgt keine Mahnung, der Fall wird zur gezielten Bearbeitung vorbereitet und der Kunde erlebt das Unternehmen als koordiniert, informiert und fair. Dieses Gedankenexperiment zeigt, dass nicht einzelne Agenten, sondern erst das Zusammenspiel der Schichten den eigentlichen Mehrwert erzeugt.

Die Einführung eines agentischen Debitorenmanagements erfordert kein radikales Big-Bang-Projekt, sondern ein strukturiertes, risikokontrolliertes Vorgehen. Der erste Schritt besteht in der vollständigen Transparenz über verfügbare Datenquellen und deren Integrationsfähigkeit. Ohne saubere, verlässliche Datenströme bleibt jede AI-Intelligenz wirkungslos.

Darauf aufbauend sollte die ereignisgesteuerte Integrationsschicht etabliert werden, um relevante Veränderungen in Echtzeit verfügbar zu machen. Parallel dazu wird der Data & Context Layer aufgebaut, der historische und aktuelle Informationen zu einem konsistenten Wissensmodell zusammenführt. Erst wenn diese Grundlage stabil ist, beginnt die eigentliche Agentenentwicklung.

In einer frühen Phase empfiehlt sich der Betrieb der Agenten im sogenannten Shadow Mode. Die AI trifft bereits Entscheidungen und formuliert Handlungsvorschläge, greift jedoch noch nicht aktiv in operative Systeme ein. Diese Phase ermöglicht den Vergleich zwischen menschlichen Entscheidungen und agentischen Empfehlungen und schafft Vertrauen in die Qualität der Logik.

Nach erfolgreicher Validierung können schrittweise autonome Schreibrechte freigegeben werden. Zunächst für risikoarme Fälle, später für komplexere Szenarien. Flankiert wird dieser Prozess durch klare Governance-Regeln, Audit-Trails und kontinuierliches Performance-Monitoring. So entsteht kein Kontrollverlust, sondern ein kontrollierter Übergang von manueller Steuerung zu intelligenter Autonomie.

Fazit

Das vorgestellte Konzept verdeutlicht, dass der Einsatz autonomer Agenten im Debitorenmanagement weit über eine reine Effizienzsteigerung hinausgeht: Es markiert einen grundlegenden Paradigmenwechsel von der reaktiven Datenverwaltung hin zur vorausschauenden, agentischen Steuerung. Durch die intelligente Verknüpfung isolierter Datenströme und deren Echtzeit-Verarbeitung entsteht ein System, das innerhalb klar definierter Leitplanken autonom agiert und Mitarbeiter konsequent von operativen Routinen entlastet. Diese gewinnen dadurch den Freiraum, sich auf ihre neue Rolle als strategische Entscheider und Gestalter der Autonomie-Matrix zu konzentrieren, während die AI die operative Geschwindigkeit sichert.



Dualer

CONSULTING

Dabei wird die technologische Innovation durch eine Architektur der Sicherheit flankiert, die Compliance und rechtliche Beherrschbarkeit ins Zentrum stellt. Indem kritische Maßnahmen konsequent an eine menschliche Letztentscheidung rückgebunden werden (Human-in-the-Loop) und jeder Agentenschritt in einem lückenlosen Audit-Trail dokumentiert wird, bleibt das System jederzeit transparent und revisionssicher. Diese harmonische Verbindung aus dynamischer Autonomie und menschlicher Kontrolle minimiert Haftungsrisiken und erfüllt bereits heute die Anforderungen künftiger AI-Regulierungen. Letztlich transformiert dieser Ansatz das Forderungsmanagement von einem fehleranfälligen Kostenfaktor zu einem digitalen, rechtssicheren Steuerungsinstrument, das in volatilen Märkten einen entscheidenden Wettbewerbsvorteil bei Liquidität und Kundenbeziehung sichert.



5.3 Rechtliche Umsetzung

Dieses Kapitel analysiert die rechtlichen Risiken der beschriebenen AI-gestützten Architektur für das Debitorenmanagement entlang ihrer technischen Ebenen. Ziel ist nicht eine vollständige Rechtsprüfung, sondern vielmehr die Identifikation der jeweils signifikantesten rechtlichen Konfliktlinien je Ebene sowie die Ableitung konkreter rechtlicher Umsetzungsmaßnahmen. Maßgeblich sind hierfür insbesondere der EU AI Act (VO (EU) 2024/1689), die DSGVO, die Grundsätze ordnungsmäßiger Buchführung (GoBD) sowie zivilrechtliche Anforderungen an die Vertragsgestaltung und das Mahnwesen.

Auf der Ebene 1 (Integration & Event Layer) erfolgt die ereignisgesteuerte Erfassung von Änderungen aus Bank, ERP und CRM in Echtzeit. Rechtlich sind hierbei insbesondere die Grundsätze der Datenverarbeitung nach Art. 5 DSGVO sowie die Rechtmäßigkeit nach Art. 6 Abs. 1 DSGVO für die automatisierte Rechnungsprüfung relevant. Da Daten aus unterschiedlichen Quellen zusammengeführt werden, besteht ein Risiko in der unzulässigen Zweckänderung der ursprünglich erhobenen Daten. Die Verarbeitung erfolgt daher ausschließlich zum Zweck der ordnungsgemäßen Vertragsdurchführung und Liquiditätssicherung, wobei der AI-Agent nur auf zwingend erforderliche Informationen zugreift.

Die Ebene 2 (Data & Context Layer) reichert isolierte Events mit der Historie und Korrespondenz zu einem konsistenten Wissensmodell an. Dieser Schritt berührt in besonderem Maße die Regelungen zur automatisierten Entscheidungsfindung gemäß Art. 22 DSGVO sowie die Transparenz- und Fairnessanforderungen. Da aus der Analyse von Zahlungshistorien Risikoscores abgeleitet werden, die wirtschaftliche Auswirkungen auf Kunden haben können, ist dieser Schritt rechtlich sensibel. Um unzulässige automatisierte Einzelentscheidungen zu vermeiden, dient die Risikobewertung ausschließlich der Entscheidungsunterstützung. Es wird sichergestellt, dass die Bewertung auf nachvollziehbaren Kriterien basiert und diskriminierende Merkmale systematisch ausgeschlossen werden.

Hinsichtlich der Ebene 3 (Agentic AI Layer), in dem spezialisierte Agenten wie der Invoice Matching oder Credit Risk Agent agieren, sind die Grundsätze der Zweckbindung und Datenminimierung nach Art. 5 DSGVO sowie zivilrechtliche Aspekte der Vertragsgestaltung relevant. Da die Agenten adaptive Vorschläge für Zahlungsziele oder Ratenpläne generieren, ist eine klare Abgrenzung zwischen automatisierter Empfehlung und verbindlicher Vertragsänderung erforderlich. Diese Vorschläge sind ausdrücklich als Entscheidungsvorlagen ausgestaltet und werden erst nach expliziter Freigabe durch einen Mitarbeiter wirksam.

Im Rahmen der Ebene 4 (Orchestration Layer) wird das Zusammenspiel der Agenten gesteuert und Zielkonflikte, etwa zwischen Liquiditätssicherung und Kundenbeziehung, aufgelöst. Rechtlich sind hier neben den Transparenzpflichten nach Art. 12-14 DSGVO insbesondere zivil- und wettbewerbsrechtliche Aspekte der Mahnsteuerung von Bedeutung. Mahnschreiben stellen rechtlich relevante Kommunikation dar und dürfen weder irreführend sein noch unbeabsichtigt rechtsverbindliche Erklärungen enthalten. Während einfache Erinnerungen automatisiert versendet werden können, ist bei sensiblen Konstellationen oder höheren Eskalationsstufen eine manuelle Prüfung vorgesehen.

Die Ebene 5 (Execution Layer) fungiert als operative Hand des Systems und transformiert digitale Logik in reale Geschäftsprozesse in den Zielsystemen. Hierbei sind die Prinzipien des Datenschutzes durch Technikgestaltung (Art. 25 DSGVO) sowie die Datensicherheit (Art. 32



DSGVO) maßgeblich. Jede Aktion, wie das Rückschreiben in das ERP-System, wird revisionssicher dokumentiert und bleibt kontrollierbar sowie gegebenenfalls reversibel. Damit wird gewährleistet, dass die operative Effizienz nicht zu einem Verlust der rechtlichen Kontrolle führt.

Zuletzt dient die Ebene 6 (Human-in-the-Loop & Governance Layer) als strategisches Sicherheitsnetz. Hier greifen die Rechenschaftspflicht gemäß Art. 5 Abs. 2 DSGVO sowie die Anforderungen der GoBD für finanzrelevante Prozesse. Der Layer implementiert Freigabe-Workflows auf Basis von Schwellenwerte, sodass kritische Maßnahmen zwingend eine menschliche Freigabe erfordern. Ein vollständiger Audit-Trail dokumentiert jede Agentenentscheidung und manuelle Eingriffe, was eine lückenlose Nachvollziehbarkeit gegenüber Wirtschaftsprüfern und Aufsichtsbehörden ermöglicht. So wird die Autonomie der AI rechtlich beherrschbar gemacht und dauerhaft in eine konforme Organisationsstruktur eingebettet.

5.4 Abschließendes Fazit zum Use Case Debitorenmanagement

Das vorgestellte Konzept verdeutlicht, dass der Einsatz autonomer Agenten im Debitorenmanagement weit über eine reine Effizienzsteigerung hinausgeht: Es markiert einen grundlegenden Paradigmenwechsel von der reaktiven Datenverwaltung hin zur vorausschauenden, agentischen Steuerung. Durch die intelligente Verknüpfung isolierter Datenströme und deren Echtzeit-Verarbeitung entsteht ein System, das innerhalb klar definierter Leitplanken autonom agiert und Mitarbeiter konsequent von operativen Routinen entlastet.

Diese gewinnen dadurch den Freiraum, sich auf ihre neue Rolle als strategische Entscheider und Gestalter der Autonomie-Matrix zu konzentrieren, während die AI die operative Geschwindigkeit sichert. Dabei wird die technologische Innovation durch eine Architektur der Sicherheit flankiert, die Compliance und rechtliche Beherrschbarkeit ins Zentrum stellt. Indem kritische Maßnahmen konsequent an eine menschliche Letztentscheidung rückgebunden werden (Human-in-the-Loop) und jeder Agentenschritt in einem lückenlosen Audit-Trail dokumentiert wird, bleibt das System jederzeit transparent und revisionssicher. Diese harmonische Verbindung aus dynamischer Autonomie und menschlicher Kontrolle minimiert Haftungsrisiken und erfüllt bereits heute die Anforderungen künftiger AI-Regulierungen. Letztlich transformiert dieser Ansatz das Forderungsmanagement von einem fehleranfälligen Kostenfaktor zu einem digitalen, rechtssicheren Steuerungsinstrument, das in volatilen Märkten einen entscheidenden Wettbewerbsvorteil bei Liquidität und Kundenbeziehung sichert.



6. Handlungsempfehlung

Die Einführung agentischer AI-Systeme sollte schrittweise und mit klar definiertem Pilotumfang erfolgen. Zunächst empfiehlt sich die Auswahl eines klar abgegrenzten Anwendungsbereichs mit überschaubarem Datenumfang und eindeutig messbaren Zielgrößen. Entscheidend ist dabei nicht die maximale technische Komplexität, sondern die saubere End-to-End-Integration entlang eines realen Geschäftsprozesses. Ein funktionsfähiger MVP mit klarer Governance, definierten KPI-Logiken und Human-in-the-Loop-Freigaben schafft früh Akzeptanz und reduziert Implementierungsrisiken.

Parallel zur technischen Umsetzung ist eine klare organisatorische Verankerung erforderlich. Rollen, Verantwortlichkeiten und Freigabelogiken müssen eindeutig definiert sein, insbesondere dort, wo Handlungsempfehlungen wirtschaftliche oder regulatorische Relevanz besitzen. Transparenzmechanismen wie Audit-Trails, versionierte Entscheidungslogiken und dokumentierte Annahmen sind von Beginn an integraler Bestandteil der Architektur. So wird sichergestellt, dass Effizienzgewinne nicht auf Kosten von Nachvollziehbarkeit oder Compliance entstehen.

Langfristig entfaltet agentische AI ihren Mehrwert erst durch kontinuierliche Weiterentwicklung und Lernen. Die Skalierung auf weitere Datenquellen, zusätzliche Agentenrollen und komplexere Entscheidungsszenarien sollte iterativ erfolgen und regelmäßig anhand klarer Erfolgskennzahlen bewertet werden. Ziel ist kein autonomes System ohne Kontrolle, sondern eine strukturierte, reproduzierbare Entscheidungsunterstützung, die Management und operative Einheiten gleichermaßen entlastet und die strategische Steuerungsfähigkeit nachhaltig stärkt.



Definitionsverzeichnis

Begriff	Definition
Abweichungskandidat (Deviation Candidate)	Vorläufige, datenbasierte Auffälligkeits-Hypothese inkl. Bewertung und Evidenz, ohne Diagnose oder automatische Maßnahme.
Fallkarte (Deviation Case)	Zentrales Arbeitsobjekt, das Muster, Kontext, begründete Priorität, Wirkungsbandbreite, Maßnahmen und Evidenz in einem steuerbaren Vorgang bündelt.
Evidenzpaket	Zusammenstellung relevanter Sensorsegmente, Ereignisse, Kontextdaten und Verweise, die die Ableitung eines Kandidaten/Empfehlung nachvollziehbar machen.
Evidenzgebundenes Arbeiten	Prinzip, dass jede Agentenaussage auf referenzierten Datenobjekten basiert; sonst als Hypothese gekennzeichnet.
Wirkungsbandbreite	Spannweite möglicher wirtschaftlicher/operativer Effekte (statt Scheingenauigkeit), inkl. Treiber und Annahmen.
Playbook	Standardisierte Vorgehensweise mit Voraussetzungen, Warnhinweisen und dokumentierter Wirksamkeit aus historischen Fällen/Outcomes.
Ähnlicher Fall / Similar Case Set	Gruppe historischer, vergleichbarer Fälle zur Begründung, Priorisierung und Maßnahmenselektion.
Operational Context	Betriebsrahmen (z. B. Schicht, Auftrag, Produkt, Betriebsmodus), der technische Muster interpretierbar macht.
Datenqualitätskennzeichen	Vertrauensmarker je Messwert (z. B. Lücken, Ausreißer, Zeitbasis), die steuern, ob ein Signal für Bewertungen zulässig ist.
Nicht-Eignungsklausel	Regel, die ungeeignete Messwerte (mangelnde Qualität) von wirkungsrelevanten Bewertungen ausschließt.
Datenvertrag	Vereinbarung, welche Felder ein Datensatz enthalten muss und welche Bedeutung sie haben - für robuste Integration/Konsistenz.
Kostenmodell-Snapshot	Versionierte Annahmenlage für monetäre Wirkungseinordnung; macht Entscheidungen review- und auditfähig.
Decision Log	Strukturierte Dokumentation der getroffenen Entscheidung (inkl. Begründung, Annahmen, Verantwortungen).
Action Outcome	Erfasste Ergebnisse einer Maßnahme (z. B. Zeit bis Stabilisierung, Wiederkehr, Qualitäts-/Ausstoßwirkung) als Lernsignal.
Statusmodell (Falllebenszyklus)	Definierte Stati eines Falls (z. B. offen, in Prüfung, mitigiert, gelöst) zur prozesssicheren Abarbeitung.
Deduplikation (Case Management)	Zusammenführung wiederholter Kandidaten desselben Musters in einen Fall anstatt Mehrfacharbeit.
Clustering (Case Management)	Gruppierung ähnlicher Fälle zur Erkennung von Wiederholungsmustern und standortübergreifenden Trends.
Kontextagent	Rekonstruiert das technische/betriebliche Gesamtbild eines Falls (Signale, Ereignisketten, Betriebsmodus).
Ähnlichkeitsagent	Findet vergleichbare historische Fälle/Playbooks und liefert referenzierte Belege.



Wirkungsagent	Übersetzt technische Abweichungen in wirtschaftliche/operative Wirkungsbandbreiten inkl. Annahmen.
Maßnahmenagent	Strukturiert nächste Schritte (Stabilisierung, Diagnose, Behebung) inkl. Voraussetzungen und Risiken.
Erkläragent (Explanation Agent)	Bereitet Ergebnisse zielgruppengerecht (Technik/Management) mit Evidenzverweisen auf.
Orchestrator (Agenten-Orchestrator)	Steuerinstanz für Reihenfolge, Vollständigkeit und Konfliktauflösung zwischen Agentenergebnissen; triggert Eskalation/Human-Review.
Versionierung (Wissen/Modelle)	Nachvollziehbare Historisierung von Kostenmodellen, Playbooks und Lernständen zur Auditierbarkeit.
Learning- & Knowledge-Layer	Schicht, die Entscheidungen, Maßnahmen und Outcomes rückkoppelt und Playbooks, Musterbibliothek und Kalibrierungen erzeugt.
Drift Monitoring / Drift-Agent	Erkennt Verschiebungen von Normalverhalten/Baselines und stößt einen kontrollierten Aktualisierungsprozess an.
Connector & Ingestion Layer	Stabile, automatisierte Datenanbindung (APIs/Exporte) mit definierten Update-Zyklen (read-only im MVP).
Data Processing Layer	Technische Bereinigung/Harmonisierung (Duplikate, Typen, Plausibilitäten) inkl. Regeln für Data Freshness.
Semantic & KPI Layer	Fachliche Harmonisierung: zentrale KPI-Definitionen, Schwellen/Ampeln, Berechnungslogiken, dokumentiert & versioniert.
Decision-Ready Layer	Verdichtung zu managementtauglichen Sichten/Decision Packs inkl. Quellen, Zeitstempeln, Definitionen, Annahmen.
Decision Pack	Standardisiertes Dossier für Entscheidungen (Status, Top-Abweichungen, Risiken, Abhängigkeiten, Optionen) mit Referenzen.
Risk Agent	Identifiziert und priorisiert Abweichungen, Trends und Risiken aus konsolidierten KPIs/Daten.
KPI-Agent	Berechnet/aktualisiert Kennzahlen entlang definierter Logiken (nicht autonom definiert).
Decision Support Agent	Erstellt entscheidungsvorbereitete Unterlagen/Optionen; trifft keine automatisierten Entscheidungen.
Guardrail/Compliance Agent	Überwacht Berechtigungen, Logging, Governance und Leitplanken (keine autonomen Off-Scope-Aktionen).
Agenten-Blueprint	Standard für Rolle, Eingaben/Outputs, Prüfungen, Grenzen, Quellenpflicht und Dokumentationsformat eines Agenten.
Prompt-Blueprint (Output-Standard)	Vorgabe, wie Agenten Ergebnisse strukturiert liefern (z. B. „Was? Warum kritisch? Bedeutung? Optionen?“) inkl. Quellen/Zeiten.
Integration & Event Layer (Debitoren)	Ereignisgesteuerte Erfassung von Änderungen (Bank, ERP, CRM, E-Mail) in Echtzeit als standardisierte Events.
Data & Context Layer (Debitoren)	Anreicherung/Einigung isolierter Events mit Historie, Korrespondenz, Vektorisierung unstrukturierter Daten zu einem konsistenten Wissensmodell.
Agentic AI Layer (Debitoren)	Modulares Set spezialisierter Agenten (z. B. Invoice Matching, Credit Risk, Payment Terms, Communication).



Orchestration Layer (Debitoren)	Koordiniert Agentenergebnisse, löst Zielkonflikte (Liquidität vs. Beziehung vs. Risiko), steuert Prozesskohärenz.
Execution Layer (Debitoren)	Operative Umsetzung mit gesicherten Schreibrechten in Zielsysteme (ERP, Mahnlauf, Tickets, Kommunikation).
Human-in-the-Loop & Governance Layer	Freigabeschwellen, Audit-Trail, Dashboards; Mensch bleibt Finalinstanz für sensible/größere Fälle.
Invoice Matching Agent	Revisionssicherer Abgleich von Zahlungen mit offenen Posten (erkennt Teil-/Überzahlung, Skonto, Reklamationsbezug); bucht nur bei hoher Sicherheit.
Credit Risk Agent	Kontinuierliche Bonitäts-/Risikobewertung (Score, Trend) auf Basis interner/externer Signale; gibt präventive Empfehlungen.
Payment Terms Agent	Risikoadaptive Zahlungsziele/Ratenpläne, dynamische Anpassung nach Zahlungsverhalten und Kundenwert.
Dunning Strategy Agent	Taktische Mahnsteuerung (Zeitpunkt, Eskalationsstufe) gemäß Risiko, Historie, Zahlungsplan, Mahnstopp.
Communication Agent	Personalisierte Kundenkommunikation (Tonalität, Kanal, Kontext) mit Wirkungshypothese; vermeidet rechtlich bindende Aussagen.
Eskalationslogik (Decision Layer)	Regelwerk, wann menschliche Freigabe verpflichtend ist (z. B. Sicherheits-/Produktionswirkung, hohe Beträge, strategische Kunden).
Feature Engineering	Ableitung aussagekräftiger Merkmale (z. B. Trendstärke, Varianz, Fensterstatistiken) aus Rohsignalen für Mustererkennung.
Feature-Katalog	Dokumentierte Liste gebildeter Merkmale mit Zweck, Herkunft, Wirkungsbereich und Abgrenzung zu bewertenden Ableitungen.
Sensor Segment	Zusammengefasstes Zeitfenster vieler Messpunkte, um technische Muster über Zeit auszuwerten.
Event Record	Ereigniseintrag (Alarm, Zustandswechsel, Qualitätsmeldung) mit Zeitstempel/Kontext.
Asset Master Objekt	Beschreibung des betroffenen Assets (Linie/Maschine/Netzelement) inkl. Kritikalität/Hierarchie.
Impact Assessment	Strukturierte Wirkungseinordnung (Bandbreite, Treiber, Annahmen) für Managemententscheidungen.
Recommendation Plan	Abgestufter Maßnahmenplan (Stabilisierung → Diagnose → Behebung) inkl. Voraussetzungen/Risiken.
Explanation Packet	Zielgruppengerechte Erklärung der Ergebnisse mit Evidenzverweisen, Unsicherheiten und offenen Punkten.
Decision Layer (Integration)	Stelle, an der priorisierte Fallkarten/Entscheidungen in den Prozess (Leitstand, Instandhaltung, Management) eingebracht werden.
Data Freshness	Regel/Prüfkriterium, das die Aktualität der Daten sicherstellt, bevor Analysen/Agentenarbeiten erfolgen.