

Von Schatten-IT zu Schatten-KI durch ChatGPT: Herausforderungen und Handlungsempfehlungen aus deutschen Unternehmen

Generative KI, allen voran Sprachmodelle wie ChatGPT, ist in vielen Unternehmen schneller angekommen als klare Regeln und Prozesse. Das Ergebnis sind einerseits produktive Abkürzungen, andererseits aber auch eine neue Welle von Schatten-IT. Mitarbeitende nutzen frei verfügbare Dienste oft ohne Freigabe, wodurch sensible Informationen in externe Systeme gelangen können. Gleichzeitig täuschen plausible, aber falsche Antworten Sicherheit vor und erschweren eine verlässliche Entscheidungsfindung. Unsere Auswertung von Praxisstimmen zeigt: Die größten Risiken liegen in Datenlecks, Halluzinationen, fehlender Transparenz, unscharfen Berechtigungen sowie offenen Fragen zu Urheberrecht und Datenschutz. Der Beitrag fasst die wichtigsten Befunde zusammen und übersetzt sie in sofort umsetzbare Maßnahmen: von Schulungen und Datenintegration über Monitoring und klare Rollenrechte bis hin zur gezielten Nutzung lizenzierter Enterprise-Lösungen oder lokaler Eigenentwicklung. Ziel ist ein pragmatischer Rahmen, der Produktivität ermöglicht und Schatten-IT wirksam begrenzt.

Malte Högemann, Christian Hauff und Oliver Thomas

Wirtschaftsinformatik & Management
<https://doi.org/10.1365/s35764-026-00588-3>
 Eingegangen: 3. Dezember 2025
 Angenommen: 23. Dezember 2025
 © The Author(s) 2026

Published online: 09 February 2026

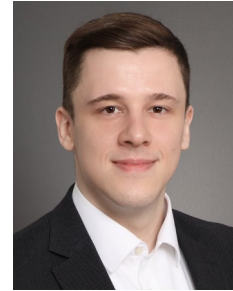
Wirtschaftsinformatik & Management

Einführung und Grundlagen

Mit der Veröffentlichung von ChatGPT im November 2022 hat generative KI einen Sprung in die Breite gemacht: Ein frei zugängliches Tool, das Texte erstellt, Code prüft, Entwürfe beschleunigt und damit den Arbeitsalltag vieler Beschäftigter prägt [1]. In verschiedenen Studien zeigte sich mittlerweile, dass generative KI im Arbeitsalltag kaum noch wegzudenken ist [2] und dabei überwiegend positive Einflüsse auf die Effizienz hat, repetitive Aufgaben erleichtert und ein hilfreiches Tool ist, um komplexe Aufgaben zu unterstützen [3].

Auf der Kehrseite zeigen sich allerdings Risiken hinsichtlich Compliance: Werden interne Inhalte in öffentliche Dienste eingegeben, drohen Datenabflüsse, das zeigen prominente Fälle bei Samsung [4] oder der australischen Regierung [5]. Zugleich verschärfen rechtliche Rahmen wie die DSGVO und der EU-AI-Act die Anforderungen an Transparenz und Verantwortlichkeiten [6, 7]. Dies gilt insbesondere bei der Nutzung von KI-Diensten außereuropäischer Anbieter, da ein Großteil der marktreifen Lösungen derzeit aus den USA bzw. anderen Nicht-EU-Ländern stammt. Dieses Spannungsfeld verstärkt sich zusätzlich, wenn KI innerhalb sensibler Branchen wie zum Beispiel der Wirtschaftsprüfung eingesetzt wird [8]. Vor diesem Hintergrund stellt sich mittlerweile nicht mehr die Frage, ob Unternehmen KI wie ChatGPT nutzen, sondern mit welchen Ansätzen aus dem Management dies sicher, transparent und wertschöpfend gelingt [9].

Unter Schatten-IT werden IT-Systeme, -Prozesse oder -Dienste verstanden, die ohne Wissen oder Freigabe der zentralen IT eingeführt und genutzt werden, um akute Bedarfslücken zu schließen [10]. Das Spektrum reicht von kleinen Tabellenlösungen bis zu integrierten Systemen oder ungeprüft beschafften Cloud-Services [11]. Die zunehmende Verfügbarkeit benutzerfreundlicher Technologien erleichtert Mitarbeitenden die eigenständige Implementierung und Nutzung von IT-Lösungen. Das wird durch die steigende technische Kompetenz der Mitarbeitenden in den Abteilungen verstärkt [10]. Generative KI und vor allem ChatGPT reißen sich nahtlos in dieses Muster ein [13]: Sie sind sofort verfügbar, können ohne Implementierungsprojekt im Browser genutzt werden und liefern unmittelbaren Nutzen, was eine Nutzung außerhalb offizieller Prozesse begünstigt. Inhalte werden „mal eben“ übersetzt, Entwürfe generiert oder Code geprüft, meist ohne Datenklassifikation, Rechteprüfung oder Freigabe. Eine aktuelle Microsoft-Studie zeigt dahin gehend: 71 % der Beschäftigten nutzen nicht genehmigte KI-Tools im Job, 51 % sogar wöchentlich. Treiber sind Bequemlichkeit und Vertrautheit



Malte Högemann^{1,2} (✉)

ist wissenschaftlicher Mitarbeiter und Doktorand am Fachgebiet Informationsmanagement und Wirtschaftsinformatik der Universität Osnabrück. Seine Forschungsschwerpunkte drehen sich vor allem darum, wie (generative) KI sicher, vertrauenswürdig und wertschöpfend sowohl in Unternehmen als auch in den Alltag von Nutzenden integriert werden kann.

malte.hoegemann@uni-osnabrueck.de



Christian Hauff¹

ist Absolvent am Lehrstuhl Informationsmanagement und Wirtschaftsinformatik der Universität Osnabrück und als Prozess- und Projektmanager im Bereich der Energiewirtschaft bei der Nowega GmbH tätig. Seine Arbeitsschwerpunkte liegen in der digitalen Transformation interner Geschäftsprozesse, mit besonderem Fokus auf der Minimierung von Schatten-IT sowie der Gewährleistung von IT-Compliance.

chauff@uni-osnabrueck.de

aus dem Privatgebrauch sowie fehlende freigegebene Alternativen im Unternehmen. Gleichzeitig unterschätzen viele das Risiko: Nur 32 % sorgen sich um den Datenschutz eingege-

bener Unternehmens-/Kundendaten, 29 % um die IT-Sicherheit. Genau diese Mischung aus hoher Nutzung und geringer Risikowahrnehmung führt zu Schatten-IT [14]. Das ist per se nicht nur „schlecht“, sondern kann unter Mitarbeitenden Innovation anstoßen und Effizienzsteigerungen bringen [12]. Ohne Leitplanken entstehen jedoch handfeste Risiken in den Bereichen:

- **Sicherheit & Datenschutz:** unkontrollierte Datenflüsse, fehlende Logs, Trainingsnutzung externer Dienste,
- **Compliance:** Verletzung regulatorischer Vorgaben (DS-GVO), Nachweislücken, uneinheitliche Behördenpraxis,
- **Qualität & Betrieb:** Halluzinationen der Modelle, fehlende Reviews, keine Dokumentation oder Tests,
- **Architektur & Kosten:** Wildwuchs an Lösungen, Medienbrüche, Mehrarbeit für Support und Pflege [10–12].

Gleichzeitig zeigen Erfahrungen aus der Praxis: Kontrollierte Schatten-IT kann Geschwindigkeit und Lernkurven erheblich verbessern, wenn die Governance mitwächst [11]. Daher stützen wir uns im Folgenden auf acht Interviews mit Experten und Expertinnen aus IT-Leitung, Informationssicherheit und Datenschutz, die bereits Erfahrungen mit der Einführung generativer KI in Unternehmen gemacht haben. Ziel ist die Ableitung praxisnaher Muster: Welche Risiken treten tatsächlich auf, welche Gegenmaßnahmen bewähren sich, und wie lassen sich Leitplanken so gestalten, dass Produktivität steigt und Schatteneffekte sinken. Die Interviews wurden entlang eines einheitlichen Leitfadens geführt, thematisch codiert und zu umsetzbaren Handlungsempfehlungen aggregiert. Dieses Vorgehen liefert eine kompakte, anwendungsorientierte Basis für die nachfolgenden Ergebnisse.

Herausforderungen

Risiken von Datenlecks

Die Interviews zeigen: Datenlecks zählen zu den größten Risiken beim Einsatz KI-basierter Systeme. Häufig speisen Mitarbeitende unbewusst sensible Informationen in öffentliche Dienste ein, etwa wenn interne Texte „mal eben“ über Übersetzungstools bearbeitet werden und so ungewollt an externe Anbieter gelangen. Eigenständig abonnierte Tools entziehen sich der Sicht der zentralen IT. Kontrolle und Überblick über Datenflüsse gehen verloren und Inhalte können in unsicheren oder unregulierten Anwendungen landen. Hinzu kommt, dass Vorfälle oft spät erkannt werden, da bestehende Sicherheitsmechanismen nicht jeden Abfluss sofort aufdecken und



Prof. Dr. Oliver Thomas^{1,2}

leitet seit 2009 das Fachgebiet Informationsmanagement und Wirtschaftsinformatik der Universität Osnabrück und ist zudem Leiter des Forschungsbereichs Smart Enterprise Engineering des Deutschen Forschungszentrums für Künstliche Intelligenz (DFKI) am Standort in Osnabrück. Schwerpunkte der angewandten Forschung sind hybride Wertschöpfung, Smart Services und digitale Transformation. Neben seiner akademischen Tätigkeit hat er die Strategion GmbH, die Didactic Innovations GmbH und die School to go GmbH gegründet und ist dort Gesellschafter. Die Unternehmen bündeln praxisorientierte Innovations- und Entwicklungskompetenzen: von der Konzeption und Umsetzung digitaler Services und smarter Geschäftsmodelle über den Transfer betriebsinformatischer Forschungsergebnisse in Unternehmen bis hin zu Lösungen für digitale Bildung und Schulentwicklung.

oliver.thomas@uni-osnabrueck.de

¹Fachgebiet Informationsmanagement und Wirtschaftsinformatik, Universität Osnabrück, Osnabrück, Deutschland

²Forschungsbereich Smart Enterprise Engineering, Deutsches Forschungszentrum für Künstliche Intelligenz GmbH, Osnabrück, Deutschland

trotz Monitoring Lücken bis zum Ereignis unbemerkt bleiben. Ein Kernproblem ist die unterschätzte Risikowahrnehmung: Viele Beschäftigte halten die Eingabe in KI-Tools für harmlos und übertragen dadurch leichtfertig vertrauliche Informationen an Dritte.

Halluzinationen

Ein wiederkehrendes Muster in den Interviews sind Halluzinationen. KI-Systeme erzeugen zwar Antworten, die plausibel klingen, aber inhaltlich falsch oder irreführend sind. Besonders kritisch ist, dass diese Fehler oft nicht sofort erkennbar sind und somit falsche Entscheidungen nach sich ziehen können. Genannt wurden Beispiele von absurden Aussagen bis hin zu fehlerhaften Angaben zu Regulierungen. Selbst wenn eigene Dokumente als Datenbasis genutzt wurden, kam es zu Fällen, bei denen Kennzahlen falsch berechnet oder erfunden wurden. Da solche Systeme Unsicherheit selten klar signalisieren, sind subtile Fehler schwer erkennbar, vor allem wenn Nutzende nicht über ausreichend Expertise verfügen.

Fehlende Transparenz

Die Intransparenz vieler KI-Systeme, sowohl bei Datenflüssen als auch bei der Entscheidungslogik, wurde als gravierend bewertet. Dadurch können Nutzende Ergebnisse nur schwer einordnen und deren Qualität nur eingeschränkt bewerten. In der Praxis bleibt oft unklar, wohin eingegebene Daten gelangen und wie sie verarbeitet werden. Zugleich sind die Entscheidungswege komplexer Modelle wie ChatGPT nicht nachvollziehbar, obwohl Datenschutzbehörden Transparenz und Erklärbarkeit einfordern. Mitarbeitende stellen daher vermehrt die Funktionsweise und Grenzen der Systeme infrage. Kontextabhängigkeiten und fehlende Prozessnachweise erschweren zusätzlich die Rückverfolgung, welcher Vorgang warum in welcher Situation ausgelöst wurde.

Berechtigungsprobleme

Unklare Rollen und Rechte stellen ein zentrales Risiko beim Einsatz von KI-Systemen wie ChatGPT oder Copilot dar. Fehlen präzise Vorgaben, kann es zu unbefugten Einblicken in sensible Informationen kommen, wenn eigene Dokumente verarbeitet werden. Historisch gewachsene IT-Landschaften sind bereits oft intransparent und vorhandene Schwachstellen in der Rechtevergabe werden durch KI weiter verschärft, da sie Daten aus unterschiedlichen Quellen kombiniert und somit Schutzgrenzen aufweicht. Zugleich greifen Systeme teils auf umfangreiche, ungeprüfte Datensätze zu, ohne dass vorab eine gezielte Filterung oder Klassifikation erfolgt. Selbst bei Enterprise-Integrationen hängt viel von der Disziplin der Nutzenden ab: Nur wenn Berechtigungen korrekt gesetzt und Ablagen sauber geführt werden, lassen sich unerlaubte Zugriffe vermeiden.

Zusammenfassung

- Generative KI steigert Tempo und Qualität spürbar, zugleich wächst das Risiko unkontrollierter Nutzung jenseits der IT-Prozesse, insbesondere in Form von Schatten-IT.
- Die Interviews aus deutschen Unternehmen zeigen als Herausforderungen Datenlecks, Halluzinationen, Intransparenz und unscharfe Berechtigungen im Spannungsfeld von DSGVO und Mitbestimmung.
- Wer klare Leitplanken, transparente Prozesse, menschliche Prüfung und Enterprise- oder eigene Lösungen etabliert, kann Produktivitätsgewinne realisieren und Sicherheitsrisiken gleichzeitig wirksam begrenzen.

Urheberrechtsfragen

Ein weiterer Befund betrifft das Urheberrecht und die Rechtklärung bei KI-Inhalten. Einerseits ist die Herkunft der Trainingsdaten oft nicht eindeutig dokumentiert, wodurch sich Unsicherheiten ergeben, ob geschützte Werke eingeflossen sind und welche Folgen das für die Nutzung hat. Andererseits ist die Rechtslage der generierten Inhalte nicht abschließend geklärt. In kreativen Prozessen stellen sich Fragen nach Urheberstellung, Nutzungsumfang und Lizenzierbarkeit, insbesondere bei Bildern oder textnahen Werken.

Datenschutz

Zentral bleibt der Datenschutz, sowohl technisch als auch organisatorisch. Die Einhaltung der DSGVO erweist sich in der Praxis jedoch als Hürde, insbesondere wenn unklar ist, wo Eingaben verarbeitet werden und ob sie zu Trainingszwecken genutzt werden dürfen. Hinzu kommen rechtliche Unsicherheiten beim Umgang mit möglichen Verstößen, vor allem wenn personenbezogene oder vertrauliche Daten in KI-Tools einfließen. Ein wiederkehrender Befund ist die fehlende Datenklassifikation: Ohne eine klare Unterscheidung zwischen sensiblen, internen und öffentlichen Daten lassen sich Schutzmaßnahmen kaum zielgerichtet umsetzen. Zudem weisen Befragte auf Bußgeldrisiken hin und fordern explizite Garantien, dass Prompts nicht zu Trainingszwecken verwendet werden.

Handlungsempfehlungen

Schulung und Sensibilisierung

Wie alle Interviews zeigen, senken gezielte Schulungen und Sensibilisierungen die Risiken beim Einsatz von KI-Syste-

men wie ChatGPT spürbar. Viele Problemfelder, von Datenlecks bis hin zu Fehleinschätzungen von KI-Ausgaben, sind wissensgetrieben. Fehlendes Verständnis für Datenregeln, Promptgestaltung und die Grenzen der Technologie verstärken Vorfälle. Unternehmen begegnen dem mit Trainings zum Thema „Welche Information darf in welches Tool?“ sowie mit kurzen Lernformaten, die im Alltag schnell greifen und nachweislich das Risiko reduzieren. Parallel dazu werden Mitarbeitende in Promptkompetenz geschult, damit Aufgaben klar strukturiert, Kontexte sauber übergeben und Qualitätskriterien formuliert werden, wodurch der Output verlässlicher wird. Für geschäftskritische Fälle wird zudem eine verbindliche menschliche Validierung eingeführt, um halluzinierte oder unvollständige Antworten abzufangen. Praxisbeispiele unterstreichen die Wirksamkeit: Auftaktveranstaltungen mit unternehmensspezifischen Anwendungsfällen schaffen Sichtbarkeit und Akzeptanz. Anstelle von pauschalen Verboten setzen einzelne Organisationen auf frühe, offene Kommunikation und klare Leitplanken, mit dem Ergebnis, dass die Schattennutzung abnimmt und die reguläre, sichere Nutzung zunimmt.

Saubere Datenbasis aufbauen

Wenn KI-Tools genutzt werden, um eigene Daten zu verarbeiten, steht und fällt die Qualität der KI-Ergebnisse mit der Qualität der Daten. Eine gut strukturierte, aktuelle und verlässliche Datenbasis ist Grundvoraussetzung für produktive Anwendungen mit ChatGPT oder Copilot. Dies zeigen die Ergebnisse der Interviews: Wenn Daten unsauber verarbeitet werden, entstehen Verzerrungen und Fehlinterpretationen, die bis hin zu klar falschen Antworten des Chatbots führen können. Entsprechend betonen die Befragten, dass sie ihre Daten vor der Nutzung durch KI bereinigen und konsistent aufbereiten (z. B. eindeutige Stammdaten, gepflegte Metadaten, Dublettenfreiheit), um Qualitätsprobleme gar nicht erst entstehen zu lassen. Ein solides Datenfundament reduziert den Aufwand für Nacharbeiten, erhöht die Verlässlichkeit der Ergebnisse und schafft die Voraussetzung, KI-Funktionen zielgerichtet und kontrolliert in Geschäftsprozesse zu integrieren.

Überwachung und Monitoring

Eine regelmäßige Überwachung ist entscheidend, um Sicherheitsrisiken, Fehlfunktionen und unerwartetes Verhalten von KI-Systemen frühzeitig zu erkennen. In der Praxis kommen Nutzungsprotokolle und Risikoanalysen zum Einsatz. An-

Kernthesen

- Mensch im Mittelpunkt: Kompetenzen, Verantwortlichkeiten und Reviews können Fehlentscheidungen trotz KI-Tempo verhindern.
- IT-Business-Alignment schlägt Verbotskultur: Partizipative Pilotierungen liefern Nutzenbelege und reduzieren Schatten-IT, wenn Fachbereiche und IT gemeinsam steuern.
- Datenklassifikation (öffentlich/intern/sensibel) und Rechtskonzepte sind das Fundament jeder KI-Nutzung.

wendungen werden nach Risikoklassen bewertet und bei Bedarf technisch gesperrt, um Schatten-IT einzudämmen und Datenverluste zu verhindern. Unternehmen, die KI in bestehende Plattformen integrieren, berichten von Vorteilen. So lasse sich Microsoft Copilot durch vorhandene Toolchains und eine durchlaufene Datenschutzfolgenabschätzung einfacher überwachen. Ergänzend wird Traffic-Monitoring genannt, um aufgerufene Dienste nachvollziehen und Auffälligkeiten zügig adressieren zu können. Wird ein Verstoß gegen Datenschutzregeln erkannt, greifen definierte Incident-Prozesse: Zunächst wird der Schweregrad bestimmt, dann werden betroffene Bereiche informiert, Meldewege geklärt und Gegenmaßnahmen eingeleitet. Wo das Monitoring eine fortbestehende unerlaubte Nutzung sichtbar macht, reagieren Unternehmen mit Zugriffssperren für offene Endpunkte und verweisen auf sichere Alternativen. Transparenz, klare Prozesse und technische Durchsetzung reduzieren Risiken und schaffen belastbare Leitplanken für den Alltagseinsatz.

Manuelle Kontrolle verankern

Trotz Automatisierung bleibt Human-in-the-Loop der zentrale Sicherheitsanker. KI-Ergebnisse müssen auf Korrektheit, Plausibilität und Vollständigkeit geprüft werden, insbesondere in geschäftskritischen Bereichen. Die Interviews betonen, dass ChatGPT eine Unterstützung bietet, aber keine Fachprüfung ersetzt. Unpräzise Prompts verstärken Fehlstellen, weshalb eine klare Aufgabenstellung und ein klarer Kontext Pflicht sind. In der Praxis zeigt sich der Trade-off: Manuelle Kontrolle kostet Zeit, erhöht aber die Entscheidungssicherheit, beispielsweise wenn automatisch erstellte Präsentationen nachgearbeitet werden müssen oder rechtliche Implikationen eine Rolle spielen. Entsprechend fordern die Befragten

verbindliche Validierungsschritte, zum Beispiel das Vieraugenprinzip mit einer Kurzcheckliste zu Quellen, Zahlen und Aussagen, um halluzinierte oder unpassende Inhalte zuverlässig abzufangen.

Selektive Nutzerwahl und Verantwortlichkeiten

Um Datenverluste zu vermeiden und die Kontrolle zu sichern, empfehlen die Interviews einen gezielten Roll-out statt einer breiten Freigabe. Den Anfang macht ein IT-affines Pilotteam, das neue Tools erprobt, Risiken versteht und belastbare Anwendungsfälle liefert. Die Lizenzen werden bedarfsgerecht vergeben, mit Blick auf Nutzen, Kosten und Risikoprofil. Bewährt hat sich zudem der Championsansatz: In den Fachbereichen benannte Key User erhalten vertiefte Schulungen und fungieren als erste Anlaufstelle. Sie verbreiten Best Practices, bündeln Fragen und halten die Nutzung in sicheren Bahnen mit dem Ergebnis von schnellerem Wissenstransfer und messbarem Nutzen bei kontrollierbarem Risiko.

Zugriffsrechte klar regeln

Die Interviews zeigen, dass die Integration von ChatGPT oder Copilot in die Unternehmens-IT nur dann sicher ist, wenn es eine klar geregelte Berechtigungskultur gibt. Least Privilege und rollenbasierte Zugriffskontrolle sollten keine Zusatzfunktion sein, sondern eine Grundvoraussetzung für eine DSGVO-konforme Trennung und nachvollziehbare Datenflüsse. In der Praxis bedeutet das: Standardzugriffe bleiben eng begrenzt und nur ausgewählte, freigegebene Dokumente sind breit sichtbar. Alles Weitere folgt dem „Need-to-know-Prinzip“. Ein Berechtigungskonzept mit regelmäßigen Rechte-Reviews verhindert, dass Mitarbeitende unbeabsichtigt Einblick in sensible Informationen erhalten oder diese in KI-Systeme weitergeben. Ein solches Szenario kann ohne klare Strukturen schnell zu erheblichem Schaden führen. Saubere Rollen, geprüfte Freigaben und dokumentierte Zugriffe sind das Fundament für jeden sicheren KI-Einsatz.

Lizenzierte und unternehmensnahe Lösungen bevorzugen

Wo immer möglich, sollten Enterprise-Lösungen oder lokale Eigenentwicklungen den Vorzug erhalten. Sie bieten Kontrollmöglichkeiten wie Logging, Datenverlustschutzfunktionen, rollenbasierte Zugriffskontrolle und klare Aussagen zur Datennutzung. In bestehende Plattformen integriert, vereinfachen sich die Datenschutzfolgenabschätzung und die operative Betreuung.

Handlungsempfehlungen

- Schulungen als Routine verankern (Do-/Don't-Beispiele, Review-Checklisten) und Verantwortliche pro Bereich benennen
- Rollenbasierte Zugriffe, Rechte-Reviews und klare Datenregeln betriebsnah umsetzen; Risiken messen, Vorfälle einordnen und Lessons Learned festhalten
- Lizenzierte, unternehmensnahe KI-Lösungen priorisieren und Nutzung selektiv ausrollen; Verbote nur zielgerichtet und befristet mit sicheren Alternativen kombinieren

Nutzungsverbote gezielt und temporär einsetzen

Mehrere Unternehmen berichten von temporären oder bereichsspezifischen Verboten für die Nutzung von ChatGPT, die vor allem auf Datenschutz- und Compliance-Bedenken zurückzuführen sind. In einzelnen Fällen wurde die Nutzung ausdrücklich untersagt, um Informations- und Datenverluste auszuschließen. Wo anfänglich weder Freigabe noch Verbot kommuniziert wurden, entstand eine Schattennutzung, die erst durch technische Sperren beendet wurde. Im Ergebnis zeigen die Interviews: Verbote können als überbrückende Maßnahme sinnvoll sein, jedoch werden sie erst in Kombination mit klaren Alternativen und transparenter Kommunikation wirksam, damit positive Effekte gestärkt werden und die Nutzung nicht in die Schatten-IT ausweicht.

Abschließende Erkenntnisse

Mit Handlungsempfehlungen und Vorgehensmustern aus der Praxis (**Abb. 1**) kann generative KI so genutzt werden, dass der Mensch im Zentrum bleibt und Schatten-KI gar nicht erst entsteht. Partizipatives Arbeiten mit Schulungen, Feedbackschleifen und klaren Leitplanken fördert Akzeptanz und Kompetenz im Alltag. Aktuelle Bitkom-Zahlen aus Deutschland unterstreichen die Notwendigkeit dieser Handlungsempfehlungen vor allem für kleinere Unternehmen: Nur 23 % der befragten Unternehmen mit unter 100 Mitarbeitenden haben bisher einen eigenen Zugang zu generativer KI bereitgestellt, bei Unternehmen mit über 500 Mitarbeitenden bereits 43 %. Gleichzeitig haben im Durchschnitt nur 23 % Regeln für den KI-Einsatz festgelegt [15]. Fachbereiche und IT sollten demnach eine ganzheitliche KI-Strategie verfolgen und die Einführung von (generativer) KI gemeinsam priorisieren, statt sie isoliert zu betrachten [8, 9]. Dabei sollten unter anderem Regeln verständ-

Abb. 1 Herausforderungen und Handlungsempfehlungen aus der Praxis


lich gemacht, Rollen und Rechte sowie Reviews sauber verankert werden, um Risiken zur Schatten-KI früh abzufedern. So wächst Vertrauen, die Produktivität steigt messbar und generative KI wird vom Risikohebel zum Produktivitätshebel.

Funding. Open Access funding enabled and organized by Projekt DEAL.

Open Access. Dieser Artikel wird unter der Creative Commons Namensnennung 4.0 International Lizenz veröffentlicht, welche die Nutzung, Vervielfältigung, Bearbeitung, Verbreitung und Wiedergabe in jeglichem Medium und Format erlaubt, sofern Sie den/die ursprünglichen Autor(en) und die Quelle ordnungsgemäß nennen, einen Link zur Creative Commons Lizenz beifügen und angeben, ob Änderungen vorgenommen wurden. Die in diesem Artikel enthaltenen Bilder und sonstiges Drittmaterial unterliegen ebenfalls der genannten Creative Commons Lizenz, sofern sich aus der Abbildungslegende nichts anderes ergibt. Sofern das betreffende Material nicht unter der genannten Creative Commons Lizenz steht und die betreffende Handlung nicht nach gesetz-

lichen Vorschriften erlaubt ist, ist für die oben aufgeführten Weiterverwendungen des Materials die Einwilligung des jeweiligen Rechteinhabers einzuholen. Weitere Details zur Lizenz entnehmen Sie bitte der Lizenzinformation auf <http://creativecommons.org/licenses/by/4.0/deed.de>.

Hinweis des Verlags

Der Verlag bleibt in Hinblick auf geografische Zuordnungen und Gebietsbezeichnungen in veröffentlichten Karten und Institutsadressen neutral.

Literatur

- [1] Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187–192. <https://doi.org/10.1126/science.adh2586>.
- [2] Singla, A., Sukharevsky, A., Yee, L., Chui, M., & Hall, B. (2025). *The state of AI: How organizations are rewiring to capture value*. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>. Accessed 18.10.2025.
- [3] Schlude A, Harles D, Stürz RA, Stumpf C (2024) Verbreitung generativer KI im privaten und beruflichen Alltag 2024. bidt – Bayerisches Forschungsinstitut für Digitale Transformation, München. <https://doi.org/10.35067/xypq-kn72>

- [4] Ray, S. (2023). *Samsung Bans ChatGPT And Other Chatbots For Employees After Sensitive Code Leak*. <https://www.forbes.com/sites/siladityaray/2023/05/02/samsung-bans-chatgpt-and-other-chatbots-for-employees-after-sensitive-code-leak/>. Accessed 18.10.2025.
- [5] Chung, F. (2025). 'Deeply sorry': NSW contractor uploaded 3000 flood victims' data to ChatGPT. <https://www.news.com.au/technology/online/internet/deeply-sorry-nsw-contractor-uploaded-3000-flood-victims-data-to-chatgpt/news-story/f8d0777ebe-a5400090be402b8314dede>. Accessed 18.10.2025.
- [6] Europäisches Parlament und Rat der Europäischen Union (2016). *Verordnung (EU) 2016/679 (Datenschutz-Grundverordnung, DSGVO)*. <https://eur-lex.europa.eu/eli/reg/2016/679/oj>. Accessed 18.10.2025.
- [7] Europäisches Parlament und Rat der Europäischen Union (2024). *Verordnung (EU) 2024/1689 des Europäischen Parlaments und des Rates vom 13. Juni 2024 über harmonisierte Vorschriften für künstliche Intelligenz und zur Änderung bestimmter Rechtsakte der Union (Gesetz über künstliche Intelligenz, „AI Act“)*. <https://eur-lex.europa.eu/eli/reg/2024/1689>. Accessed 18.10.2025.
- [8] Thomas, O. (2025): Künstliche Intelligenz im Finanz- und Rechnungswesen: Ein strukturierter Ansatz und Implikationen für die Wirtschaftsprüfung. In: *iwp-Journal*, Ausgabe 2/2025, 65–76.
- [9] Thomas, O. (2024): KI-Management – Ein ganzheitlicher Ansatz für die erfolgreiche Nutzung und Implementierung von Künstlicher Intelligenz in Unternehmen. In: Thomas, O. (Hrsg.): *Strategien Report*, Nr. 1, Osnabrück, Strategien GmbH
- [10] Zimmermann, S., & Rentrop, C. (2012). Schatten-IT. *HMD – Praxis der Wirtschaftsinformatik*, 49(6), 60–68. <https://doi.org/10.1007/BF03340758>.
- [11] Rentrop, C., Zimmermann, S., & Huber, M. (2015). Schatten-IT – ein unterschätztes Risiko. In P. Schartner, K. Lemke-Rust & M. Ullmann (Eds.), *D·A·CH Security 2015* (pp. 291–300). syssec.
- [12] Fürstenau, Daniel; Rothe, Hannes. Shadow IT Systems: Discerning the Good and the Evil. In: *Proceedings of the 22nd European Conference on Information Systems (ECIS 2014)*, Tel Aviv, 9.–11. Juni 2014. AIS Electronic Library (AISeL).
- [13] Silic, M., Silic, D., & Kind-Trüller, K. (2025). From Shadow IT to Shadow AI—Threats, Risks and Opportunities for Organizations. *Strategic Change*.
- [14] Microsoft (2025). *Rise in 'Shadow AI' tools raising security concerns for UK organisations*. <https://ukstories.microsoft.com/features/rise-in-shadow-ai-tools-raising-security-concerns-for-uk/>. Accessed 18.10.2025.
- [15] Bitkom e. V. (2025). *Beschäftigte nutzen vermehrt Schatten-KI*. <https://www.bitkom.org/Presse/Presseinformation/Beschaeftigte-nutzen-Schatten-KI>. Accessed 22.10.2025.