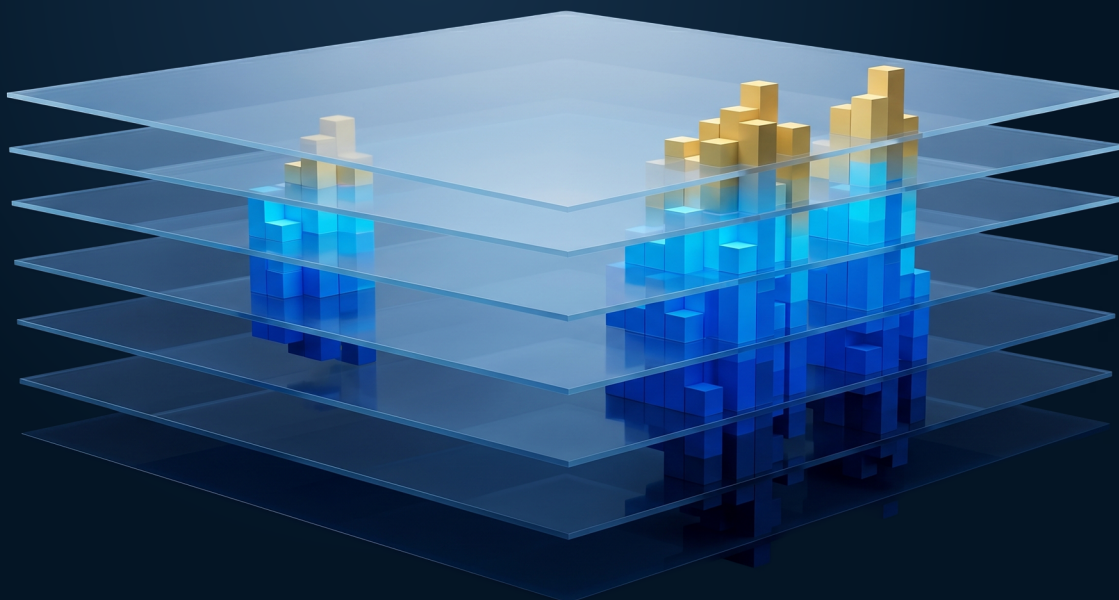




Bayerischer
KI-Innovationsbeschleuniger

Wer nicht wagt, der nicht gewinnt?

Empfehlungen für beschleunigte
Konformität mit der EU KI-Verordnung
basierend auf der Risikoklassifizierung
von 628 deutschen KI-Startups



 **appliedAI**
institute
for europe



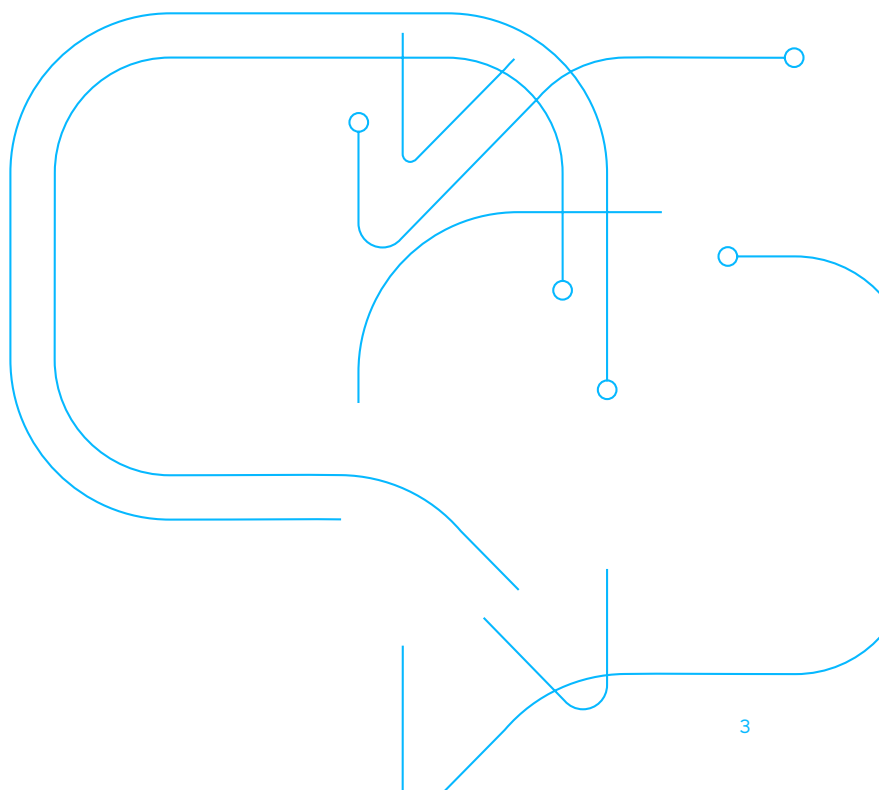
Gefördert durch
Bayerisches Staatsministerium
für Digitales



Inhaltsverzeichnis

	Vorwort	4
	Kurzfassung	5
1.	Motivation	8
2.	Warum die Risikoklassifizierung gemäß KI-Verordnung für KI-Startups von Bedeutung ist	9
	2.1 Allgemeine Risikoklassifizierung nach der KI-Verordnung	10
	2.2 Die Kosten eines hohen Risikos	10
	2.3 Unterstützung für KI-Startups	11
3.	Ziele, Methoden und Daten	12
	3.1 Ziele der Studie	12
	3.2 Methoden und Daten	12
4.	Ergebnisse & Beobachtungen	13
	Ergebnis 1: Der Anteil an KI-Startups mit hohem Risiko liegt deutlich über den ursprünglichen Annahmen der Europäischen Kommission	14
	Ergebnis 2: Kontext und fehlende Details schränken die vollautomatische Risikoklassifizierung ein und erschweren die manuelle Überprüfung	16
	Ergebnis 3: Die Verbreitung von GPAI-Modellen macht die Transparenzpflichten zur vorherrschenden regulatorischen Grundlage für KI-Startups	19
	Ergebnis 4: Die Compliance-Belastung betrifft bestimmte Branchen und Unternehmensfunktionen überproportional	21
	Ergebnis 5: Deutsche Bundesländer weisen ähnliche Profile hinsichtlich der Risikoklasse auf, unterscheiden sich jedoch in ihrer Verteilung auf die einzelnen Branchen	23
	Ergebnis 6: Compliance-Kosten spiegeln sich (noch) nicht in ihrer Finanzierung wider	26
5.	Schlussfolgerung	29

Anhang	30
Anhang 1: Datensatz	30
Anhang 2: KI-Tool zur Risikoklassifizierung	31
Referenzen	37
Über die Autoren	38
appliedAI Institute for Europe – Über uns	40



Vorwort

Mit dem Inkrafttreten der KI-Verordnung am 1. August 2024 stehen europäische KI-Startups vor einer entscheidenden Herausforderung: Sie müssen umfangreiche Compliance-Verpflichtungen erfüllen und gleichzeitig mit Ressourcenengpässen und Wettbewerbsdruck umgehen. Die KI-Verordnung als Rechtsrahmen spielt eine wesentliche Rolle für die Entwicklung vertrauenswürdiger KI, jedoch stellt sie eine erhebliche Belastung für Startups dar, die ohnehin schon mit begrenzten finanziellen und personellen Ressourcen zu kämpfen haben.

Am appliedAI Institute for Europe erleben wir aus erster Hand, wie deutsche KI-Startups einen essentiellen Beitrag zur Stärkung der Innovationskraft Europas leisten. Die durch die KI-Verordnung eingeführte zusätzliche regulatorische Komplexität könnte allerdings sowohl ihr Wachstum als auch ihre Skalierbarkeit erschweren. Diese Studie liefert eine umfassende, datengestützte Analyse der regulatorischen Anforderungen, die Startups betreffen werden, und bietet konkrete Empfehlungen zur Gestaltung wirksamer Unterstützungsmechanismen.

Unser Ziel der Analyse der Risikoklassifizierung der deutschen KI-Startup-Landschaft auf Grundlage der KI-Verordnung ist, diese Lücke in der bestehenden Forschung zu adressieren. Wir wollen dazu beitragen, strategische Entscheidungen datenbasiert treffen zu können und die Einhaltung von Vorschriften mit Innovation in Einklang zu bringen. Dieser Bericht bildet einen wichtigen Baustein, um sicherzustellen, dass die EU ihre weltweite Rolle im Bereich der KI weiter ausbaut und gleichzeitig ein Umfeld fördert, in dem Startups wachsen und gedeihen können.

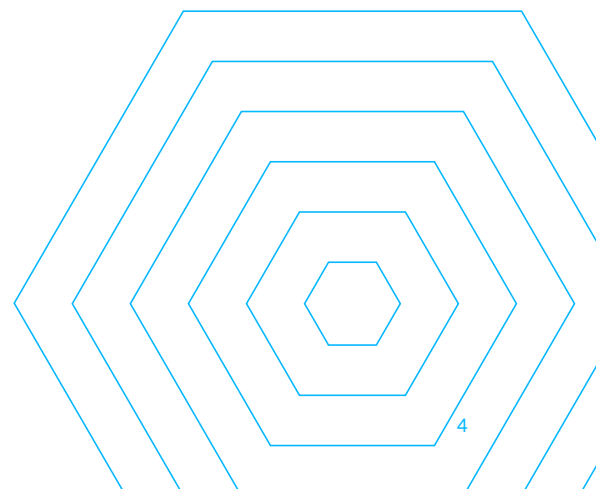
Diese Studie wurde im Rahmen des „Bayerischen KI-Innovationsbeschleunigers“ durchgeführt, der vom Bayerischen Staatsministerium für Digitales finanziert wird. Das Projekt zielt darauf ab KMU, Startups und den öffentlichen Sektor bei der Einhaltung des EU-KI-Gesetzes zu unterstützen. In diesem Kontext leistet diese Studie einen wichtigen Beitrag für einen informierten Dialog zwischen diesen Akteuren und gibt Empfehlungen für eine Zukunft, in der Innovation und Konformität Hand in Hand gehen.

Mit freundlichen Grüßen,



Dr. Wolfgang Fischer

Director Operations, Projects & Policy
appliedAI Institute for Europe

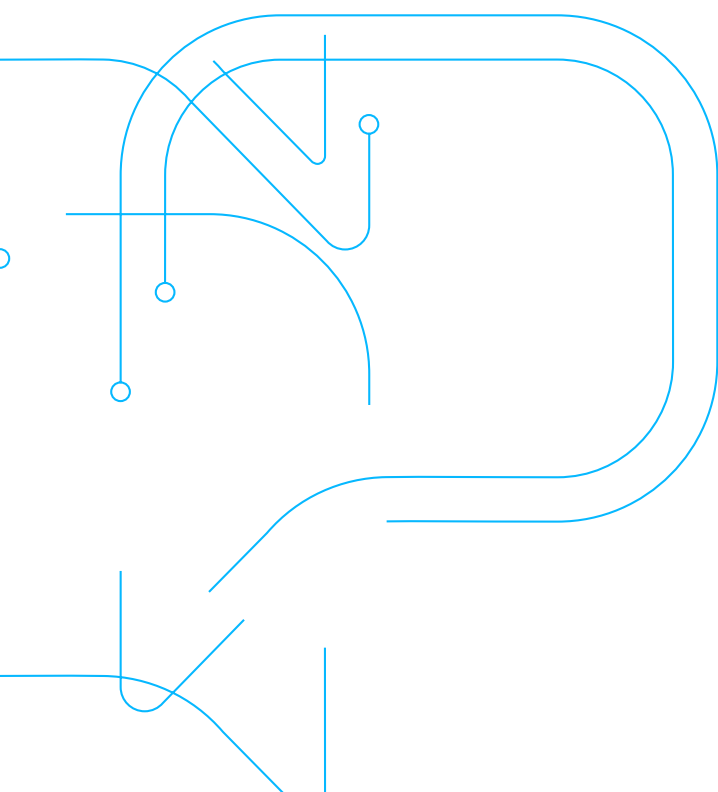


Kurzfassung

Die EU-KI-VO führt ein abgestuftes Regulierungssystem für KI-Systeme mit vier Stufen ein: KI-Systeme mit minimalem oder keinem Risiko, mit Transparenzpflichten, mit hohem Risiko sowie verbotene KI-Systeme. Die Verteilung der KI-Anwendungsfälle in Deutschland auf diese Risikostufen lässt sich derzeit jedoch noch nicht verlässlich quantifizieren.

Diese Studie nutzt einen neuartigen LLM-basierten Ansatz, um die Anwendungsfälle mit dem höchsten Risiko von 628 deutschen KI-Startups gemäß der EU KI-VO zu klassifizieren. Alle Risikoklassen wurden zusätzlich manuell überprüft und bestätigt. Die Auswertung ermittelt, wo sich der Regulierungsaufwand voraussichtlich in Hinsicht auf Risikostufen, Branchen und Unternehmensfunktionen, Bundesländer sowie Finanzierungsprofile konzentrieren wird.

Seit Inkrafttreten der KI-Verordnung ist diese Studie die bislang umfangreichste ihrer Art zur Klassifizierung von KI-Systemen. Sie ist als Vorstudie konzipiert, um weitere Forschungen zu den regulatorischen Auswirkungen der KI-Verordnung und den damit verbundenen Compliance-Kosten für das deutsche Innovations-Ökosystem anzuregen.



Wichtigste Ergebnisse

- **Die Zahl der risikoreichen Startups ist höher als allgemein angenommen:** Nach Ausschluss der Fälle, die nicht eindeutig klassifiziert werden konnten (69 von 628, d. h. 559), wird fast die Hälfte der Startups als solche mit minimalem Risiko eingestuft (48,3 % = 270/559), während etwa ein Viertel Transparenzpflichten unterliegt (26,3 % = 147/559) und etwa ein Viertel als hochriskant gilt (25,4 % = 142/559). Keiner der untersuchten Anwendungsfälle wurde als verboten eingestuft.
- **Die Risikoklassifizierung ist häufig kontextabhängig:** 69 Startups konnten aufgrund unvollständiger Informationen nicht eindeutig klassifiziert werden oder weil die Klassifizierung vom Einsatzkontext und von technischen Details abhängt, die nicht öffentlich verfügbar sind.
- **Transparenzpflichten werden sich voraussichtlich als neuer Standard etablieren:** Die Studie analysiert eine Teilmenge von 111 Startups im Bereich der Generativen KI und stellt fest, dass 59 % von ihnen Transparenzpflichten und 12 % Pflichten für Hochrisiko-KI-Systeme haben. Dies deutet darauf hin, dass sich Transparenzanforderungen mit der zunehmenden Verbreitung von LLM-basierten Lösungen zu einer vorherrschenden „Grundanforderung“ entwickeln könnten und dass das Zusammenspiel zwischen den Verhaltenskodizes (Codes of Practice) für GPAI-Modelle und den Anforderungen an Hochrisiko-KI-Systeme die nachgelagerte Compliance erheblich beeinflussen wird.
- **Der Compliance-Aufwand betrifft bestimmte Sektoren und Unternehmensfunktionen in besonderem Maße:** Startups mit hohem Risiko sind überproportional stark in Sektoren wie „Gesundheits- und Sozialwesen“ und „Verkehr“ sowie in Unternehmensfunktionen wie „Personalwesen“ und „Betrieb“ vertreten. Diese Konzentration impliziert, dass Compliance-Unterstützungsmaßnahmen, die auf den jeweiligen Sektor und die Unternehmensfunktion zugeschnitten sind, wahrscheinlich eine größere Wirkung erzielen werden.
- **Die Größe eines regionalen Ökosystems verändert die Zusammensetzung der Risikostufen insgesamt nicht wesentlich, aber Unterschiede bei den ansässigen Branchen erzeugen einen unterschiedlichen Bedarf an KI-Governance:** Führende Startup-Bundesländer weisen weitgehend ähnliche proportionale Risikoprofile auf (z. B. Berlin und Bayern), doch die Branchenzusammensetzung unterscheidet sich je nach Bundesland, woraus sich unterschiedliche Anforderungen an die fachliche Expertise der Aufsichtsbehörden ergeben.
- **Die derzeitige Finanzierung spiegelt die Compliance-Kosten noch nicht wider:** Obwohl ein Viertel aller Systeme voraussichtlich ein hohes Risiko aufweist, lassen die privaten Finanzierungsmuster noch keine Differenzierung zwischen den Risikostufen erkennen. Dies ist jedoch eine rein deskriptive Feststellung und spiegelt nicht zwangsläufig eine mangelnde Bereitschaft der Investoren wider.

Empfehlungen

Adressat	Empfehlung
EU / AI Office	• Annahmen aktualisieren und Umsetzung unterstützen: Annahmen zu den Auswirkungen auf den Anteil an Systemen mit hohem Risiko sollten überprüft und entsprechende Anpassungen der unterstützenden Maßnahmen, einschließlich der Finanzierungsmechanismen, erwogen werden. • Aufschub der Verpflichtungen für Systeme mit hohen Risiken und Transparenzpflichten: Die Durchsetzung der im Omnibus-Paket vorgeschlagenen Verpflichtungen in Bezug auf Systeme mit hohen Risiken und Transparenzpflichten sollte aufgeschoben werden, damit Unternehmen über die Ressourcen und die Zeit verfügen, die technischen Standards und den Verhaltenskodex für transparente Systeme umzusetzen.
Bundesregierung / Bundesländer	• Aufbau von Aufsichtskapazitäten, wo sich die Belastung konzentriert: Ausstattung der zuständigen Behörden und Aufsichtsbehörden mit ausreichendem technischem und rechtlichem Fachwissen, um mit einer höheren als erwarteten Anzahl an regulierten Systemen umgehen zu können. • Praktische Leitlinien beschleunigen: Aufsichtsbehörden mit einem Mandat zur Förderung von KI-Innovationen ausstatten; Maßnahmen von KI-Anbietern und Betreiber nach "besten Kräften" anerkennen und Unterstützung anbieten. Strafen und Bußgelder sollten das letzte Mittel sein. • Länderübergreifende Koordination bei branchenspezifischen Anwendungsfällen: Einrichtung länderübergreifender Mechanismen zum Austausch von Fachwissen und zur Angleichung der Auslegung für die am stärksten betroffenen Branchen und Unternehmensfunktionen.
Industrie / Startups und Investoren	• Mit der Risikoklassifizierung als Lebenszyklusprozess umgehen: KI-Startups sollten bereits in einer frühen Phase ihrer Produktentwicklung eine Risikoklassifizierung vornehmen und sich bewusst sein, dass Produktänderungen und Einsatzkontexte die Risikostufe verändern können, insbesondere da viele Startups branchen- und unternehmensfunktionsübergreifend tätig sind. • Berücksichtigung regulatorischer Pflichten und Compliance-Mindset bei der Finanzierung: Investoren sollten Compliance-Kosten und für die Umsetzungsbereitschaft in ihrer Due Diligence und Bewertung von Startups einbeziehen.

Wesentliche Limitationen der Studie

- **Neuartiger methodischer Ansatz:** Diese Klassifizierung kombiniert automatisierte LLM-Analysen mit menschlicher Überprüfung. Dies ermöglicht zwar Skalierbarkeit, aber beide Verfahren weisen eine gewisse Fehlerquote auf.
- **Fehlende Leitlinien und fehlerhafte Klassifizierungen:** Innerhalb der KI-Verordnung existieren nach wie vor Auslegungsspielräume, bis weitere Leitlinien, insbesondere zur Identifizierung von Anwendungsfällen mit hohem Risiko, von der Kommission finalisiert werden. Es ist möglich, dass menschliche Prüferinnen und Prüfer KI-Anwendungsfälle falsch klassifiziert haben.

1. Motivation

Vor der Einführung der EU KI-VO ging die Europäische Kommission in ihrem Impact Assessment von 2021 davon aus, dass die meisten KI-Anwendungsfälle in die Kategorie „Minimales oder kein Risiko“ fallen würden, während „schätzungsweise nur 5 % bis 15 % aller KI-Anwendungen ein hohes Risiko darstellen“ (Europäische Kommission, 2021, S. 69 (übersetzt))¹. Seit der Veröffentlichung der KI-Verordnung gab es jedoch keine systematischen Bemühungen, diese Erwartung zu bewerten oder zu validieren².

Ein höherer Anteil an Hochrisiko-KI-Systemen oder solchen mit Transparenzpflichten würde zu höheren Compliance-Kosten für europäische Unternehmen führen. Dies gilt insbesondere für Startups, die wahrscheinlich überproportional von den im Rahmen der EU KI-VO eingeführten Regulierungs- und Compliance-Verpflichtungen betroffen sein werden.

So hat beispielsweise ein aktueller OECD-Bericht zu den Rahmenbedingungen für KI in Deutschland hervorgehoben, dass deutsche KI-Startups auf dem globalen Markt strukturell benachteiligt sind und insbesondere im Vergleich zu den Vereinigten Staaten weniger Zugang zu Risikokapital, Tech-Fachkräften und Rechenleistung haben (OECD, 2024). Compliance-Aufgaben werden nun zusätzliche Ressourcen erfordern, die KI-Startups bereits jetzt nur schwer aufbringen können. Derzeit ist zudem zu beobachten, dass Gründerinnen und Gründer von Startups mit generativer KI die negativen Auswirkungen der EU KI-VO durch zusätzlich benötigte Ressourcen und hohe regulatorische Komplexität hervorheben (Hutchinson *et al.*, 2024).

Unsere Studie zielt darauf ab, eine strukturierte und datengestützte Bewertung der Risikoklassenverteilung deutscher KI-Startups im Rahmen der KI-Verordnung vorzunehmen. Wir kombinieren diese Bewertung mit Daten über die branchenspezifische, geschäftliche und geografische Tätigkeit der KI-Startups in Deutschland, um die regulatorischen Konsequenzen der KI-Verordnung für diese Akteure besser zu verstehen³.

Wir hoffen, dass empirische Erkenntnisse über die Verteilung der Risikoklassifizierungen im deutschen KI-Startup-Ökosystem nicht nur die Transparenz erhöhen, sondern auch die Grundlage für weitere Forschung, fundiertere datengestützte politische Handlungsempfehlungen und strategische Maßnahmen bilden. Dies kann Deutschlands Position als Vorreiter bei KI-getriebener Innovation im Einklang mit europäischen Werten stärken.

1 Weitere Einzelheiten zur Risikoklassifizierung finden sich in Kapitel 2.

2 Andere Studien führten Risikoklassifizierungen auf der Grundlage früherer Fassungen der EU KI-VO durch, das sich hinsichtlich der Kriterien und Regeln für die Risikoklassifizierung erheblich von der aktuellen Gesetzgebung unterscheidet. Weitere Einblicke in die Analyse früherer Fassungen der EU KI-VO finden sich bei Hauer *et al.* (2023) und appliedAI (2023).

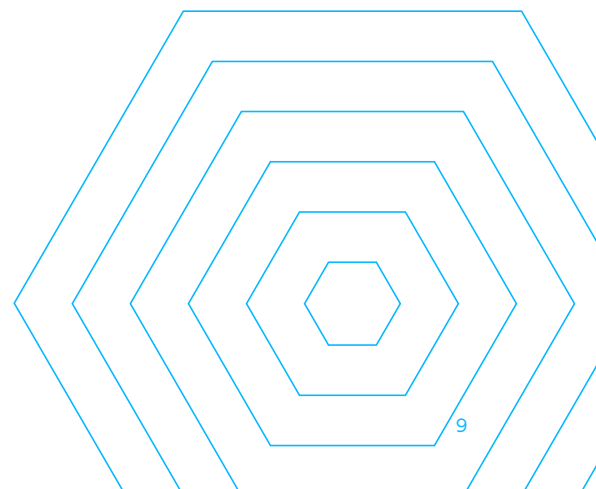
3 Diese Studie wurde vor der Veröffentlichung des Entwurfs der Leitlinien der Europäischen Kommission zur Einstufung als Hochrisikosystemen gemäß Artikel 6 (veröffentlicht am 19. Mai 2026) sowie vor der vorläufigen Einigung über das „Digital Omnibus on AI“ (7. Mai 2026) fertiggestellt, durch die die Verpflichtungen für Hochrisikosysteme auf den 2. Dezember 2027 (eigenständige Systeme) bzw. den 2. August 2028 (eingebettete Systeme) verschoben werden. Unsere Klassifizierung spiegelt den zum Zeitpunkt der Analyse geltenden Rahmen wider. Da der Omnibus-Entwurf eher die Fristen als die materiellen Einstufungsregeln anpasst, erwarten wir keine wesentlichen Auswirkungen auf die dargestellte Verteilung der Risikoklassen.

Warum die Risikoklassifizierung gemäß KI-Verordnung für KI-Startups von Bedeutung ist

Die Pflichten für einen Akteur gemäß der KI-Verordnung hängen von zwei wesentlichen Faktoren ab: der Risikoklasse seines KI-Systems und seiner Rolle in Bezug auf dieses System.

Im Rahmen der KI-Verordnung unterliegen die Anbieter von KI-Systemen – d.h. diejenigen, die diese entwickeln und auf den Markt bringen – wesentlich strengeren Auflagen als Betreiber – d.h. diejenigen, die diese lediglich nutzen. Neben der Rollenbestimmung müssen Akteure ihr KI-System zudem in das vierstufige Klassifizierungssystem der KI-Verordnung einordnen. Das zugrundeliegende Ziel dieser Systematik besteht darin, sicherzustellen, dass die Verpflichtungen in einem angemessenen Verhältnis zum Risiko stehen – je höher das Risiko einer bestimmten KI-Anwendung ist, desto größer ist die Notwendigkeit für eine risikominimierende Entwicklung.

Unsere Beobachtungen stützen die Ansicht, dass die Mehrheit der KI-Startups voraussichtlich Anbieter im Sinne der KI-Verordnung sein wird. Je höher der Anteil dieser Startups, die auch in die Hochrisiko-KI-Klasse fallen, desto höher sind die Compliance-Kosten, mit denen das Ökosystem konfrontiert ist. Dementsprechend skizziert dieses Kapitel das allgemeine System der Risikoklassifizierung gemäß der KI-Verordnung, untersucht dessen finanzielle Folgen und analysiert, was dies konkret für KI-Startups in Deutschland bedeutet.



2.1 Allgemeine Risikoklassifizierung nach der KI-Verordnung

Generell sind nach der EU KI-VO alle Akteure im Bereich von KI-Systemen verpflichtet, ihre Systeme nach einer vierstufigen Risikosystematik einzustufen. Diese Einstufung basiert auf dem Risiko, das ein KI-System für Gesundheit, Sicherheit und Grundrechte darstellt, bestimmt durch seinen Verwendungszweck.

Der Rechtsrahmen legt folgende Kategorien fest:

Unannehmbarem Risiko

KI-Systeme **mit unannehmbarem Risiko** sind gemäß Artikel 5 der KI-Verordnung verboten. Dabei handelt es sich um Systeme, die eine eindeutige Gefahr für Einzelpersonen oder die Gesellschaft darstellen, wie beispielsweise solche, die durch unterschwellige Techniken das Verhalten manipulieren, die Schutzbedürftigkeit bestimmter Bevölkerungsgruppen ausnutzen oder staatlich betriebene soziale Bewertungssysteme ermöglichen.

Hochrisiko

Hochrisiko-KI-Systeme sind zulässig, unterliegen jedoch strengen regulatorischen Anforderungen. Gemäß Artikel 6 der KI-Verordnung gibt es zwei Wege zu dieser Einstufung: Systeme, die Teil von Produkten sind, die unter bestehende EU-Produktsicherheitsgesetze fallen (aufgeführt in **Anhang I**), und eigenständige Systeme, die in sensiblen Bereichen eingesetzt werden, die in **Anhang III** aufgeführt sind, wie beispielsweise Bildung, Beschäftigung, Strafverfolgung und kritische Infrastruktur.

Transparenzpflichten

Bestimmte KI-Systeme unterliegen gemäß Artikel 50 der KI-Verordnung besonderen **Transparenzpflichten**. So müssen Nutzer beispielsweise klar darüber informiert werden, wenn sie mit KI interagieren, wie z. B. bei Chatbots. Zudem müssen synthetisch erzeugte Inhalte mit einem Wasserzeichen versehen werden.

Minimalem oder keinem Risiko

KI-Systeme mit **minimalem oder keinem Risiko**, d.h. alle Systeme, die nicht unter eine der drei vorgenannten Kategorien fallen, sind im Rahmen der KI-Verordnung weitgehend unreguliert. Zwar gelten keine verbindlichen Anforderungen, doch werden Anbieter gemäß Artikel 95 der KI-Verordnung dazu angehalten, freiwillig bewährte Verfahren anzuwenden.

2.2 Die Kosten eines hohen Risikos

Jede Rechtsvorschrift ist naturgemäß mit Kosten verbunden. Startups, die Hochrisiko-KI-Systeme entwickeln, unterliegen im Vergleich zu ihren Pendants mit minimalem oder keinem Risiko einem breiteren Spektrum strenger Anforderungen. Studien, die die Folgenabschätzung der Europäischen Kommission im Jahr 2021 stützten, schätzten, dass „...für ein Unternehmen mit 50 Mitarbeitenden die Gesamtkosten für die Einrichtung und Aufrechterhaltung der Compliance mit der KI-Verordnung im ersten Jahr für ein Produkt zwischen 216.000 EUR und 319.000 EUR liegen.“ (Renda *et al.*, 2021, S. 152–153 (übersetzt)).

Eine indikative Abschätzung der Kosten für die Einhaltung der KI-Verordnung ist zwar schwierig (Laurer *et al.*, 2021; Mueller, 2021) und hängt maßgeblich vom Konformitätsbewertungsverfahren ab, aber Startups sehen sich oft aus strukturellen Gründen mit hohen Kosten konfrontiert. Zum einen lassen sich Compliancemaßnahmen oft nicht rein intern umsetzen und erfordern möglicherweise externe Berater oder Rechtsanwälte. Zudem könnten Startups, die bisher in der Regel nicht in regulierten Umgebungen tätig waren, mit einer weitaus steileren Lernkurve konfrontiert sein, was eine Weiterqualifizierung und zusätzliches Personal erfordert, um die für die Einhaltung der Vorschriften notwendigen neuen Kompetenzen zu entwickeln. Darüber hinaus verfügen Startups möglicherweise nicht über das Know-how oder die finanziellen Ressourcen, um die Zusammenarbeit mit benannten Stellen und den Zertifizierungsprozess zu bewältigen. Startups werden mit diesen initialen und laufenden Kosten konfrontiert, während sie gleichzeitig stark auf externe Finanzmittel angewiesen sind, um ihre Geschäftsmodelle zu validieren und zu skalieren.

2.3 Unterstützung für KI-Startups

Zwar müssen alle Akteure die gleichen Verpflichtungen gemäß der KI-Verordnung erfüllen, doch gibt es Sonderbestimmungen, die speziell auf **KMU und Startups** zugeschnitten sind. In Anerkennung der wichtigen Rolle, die diese Akteure für die Innovation spielen, verpflichtet die KI-Verordnung die EU-Mitgliedstaaten, die Europäische Kommission und weitere Stellen, ihnen rechtliche und technische Unterstützung zu gewähren.

Diese Bestimmungen lassen sich grob in vier Kategorien einteilen (weitere Einzelheiten finden sich in Kapitel VI der KI-Verordnung):

- **Informationsverbreitung:** Die EU-Mitgliedstaaten sind verpflichtet, maßgeschneiderte Schulungsmaßnahmen für KMU und Startups zu organisieren und spezifische Kanäle einzurichten, um Fragen zum Gesetz und dessen Umsetzung zu beantworten.
- **Senkung der Compliance-Kosten:** KMU und Startups profitieren von weniger strengen Anforderungen an die technische Dokumentation sowie von geringeren Gebühren der zuständigen Behörden während des Compliance-Prozesses.
- **Beteiligung an Entscheidungsprozessen:** Die Mitgliedstaaten und die EU-Kommission sind verpflichtet, die Beteiligung von KMU und Startups an Entscheidungsgremien wie dem Beratungsforum und den Normungsinstituten zu erleichtern.
- **Vorrangiger Zugang zu KI-Reallaboren:** Startups und KMU profitieren von einem leichteren Zugang zu KI-Reallaboren, d.h. kontrollierten Umgebungen, in denen Unternehmen mit Produkten experimentieren können, deren regulatorischer Status ungewiss ist⁴.

Während der Großteil der KI-Verordnung im August 2026 in Kraft treten wird, gibt es kein einheitliches Datum für das Inkrafttreten all dieser Maßnahmen. Es ist zudem unklar, inwieweit sich diese Maßnahmen für Startups als kostensparend erweisen werden, da Unsicherheit in Bezug auf die tatsächlichen Kosten für die Konformität und die spezifischen Ausnahmeregelungen für kleine Anbieter besteht. KMU und Startups sollten Orientierungshilfen der Mitgliedsstaaten und lokaler Startup-Verbände heranziehen, um zu verstehen, wie sie diese Unterstützung nutzen können.

Vor diesem Hintergrund ist eine strukturierte und datengestützte Analyse der Verteilung deutscher KI-Startups auf diese Risikokategorien unverzichtbar. Nur wenn regulatorische Diskussionen auf empirischen Erkenntnissen basieren, können politische Entscheidungsträger und Branchenakteure gezielte Unterstützungsmaßnahmen entwerfen, welche Compliance und Innovation in Einklang bringen.

4 Weitere detaillierte Informationen zu regulatorischen Sandboxes, einschließlich ihrer Relevanz für KI-Startups, finden sich in Fetic *et al.* (2025).

3.1 Ziele der Studie

Während ein von der Europäischen Kommission durchgeführtes Impact Assessment einen Anteil von 5–15 % an Hochrisiko-KI-Systemen schätzt, existieren für das deutsche Startup-Ökosystem keinerlei Schätzungen bzw. Kennzahlen zur Verteilung der Risikoklassifizierungen.

Dementsprechend **ist das zentrale Ziel unserer Studie:**

Wie verteilt sich die deutsche KI-Startup-Landschaft auf die Risikostufen der EU KI-VO? Wie stellt sich dies im Vergleich zu früheren politischen Annahmen dar?

Auf der Grundlage zusätzlicher Daten, die uns aus unseren Daten zur deutschen KI-Startup-Landschaft zur Verfügung stehen, analysieren wir außerdem:

- **Branchen:** In welchen Branchen und Unternehmensfunktionen sind Hochrisiko-KI-Systeme und Transparenzpflichten überproportional vertreten?
- **Bundesländer:** Wie variieren die Verteilung der Risikostufen und die Branchenzusammensetzung in den einzelnen deutschen Bundesländern?
- **Finanzierung:** Wie sehen mögliche Finanzierungsmuster für Startups in den verschiedenen Risikoklassen aus?

3.2 Methoden und Daten

Wir untersuchen das regulatorische Umfeld deutscher KI-Startups im Rahmen der EU KI-VO, um ein umfassendes Verständnis für die Verteilung auf die verschiedenen Risikokategorien sowie die Auswirkungen auf Compliance und Aufsichtsbehörden zu ermitteln. Hierzu analysieren wir, wie sich die Risikoklassifizierungen je nach Technologieart, KI-Fähigkeiten und Unternehmensfunktionen unterscheiden. Dabei werden auch geografische und finanzielle Aspekte berücksichtigt, um Muster in der Konzentration von Startups mit Hochrisiko-KI-Systemen oder Transparenzpflichten aufzuzeigen und zu verdeutlichen, wie sich die Investitionsdynamik mit den regulatorischen Anforderungen überschneidet.

Dafür haben wir einen Datensatz mit 628 deutschen KI-Startups und 2.451 KI-Anwendungsfällen zusammengestellt, der auf öffentlich zugänglichen Beschreibungen (Websites und Produktmaterialien) basiert. Jeder Anwendungsfall wurde unter Verwendung einer LLM-gestützten Klassifizierungspipeline gemäß der Risikoklassifizierung der EU KI-VO klassifiziert, gefolgt von einer indikativen juristischen Bewertung durch unsere Expertinnen und Experten mittels einer manuellen Überprüfung. Für jedes Startup geben wir die risikoreichste Klassifizierung innerhalb seines Portfolios an, um das Compliance-Risiko im “schlimmsten Fall” widerzuspiegeln. Zusätzliche methodische Erläuterungen, einschließlich der Erstellung des Datensatzes, der Modell-/Prompt-Einrichtung, der Entscheidungsregeln für die Klassifizierung und der Validierungsergebnisse, befinden sich in den Anhängen I und II.

Ergebnisse & Beobachtungen

Dieser Abschnitt präsentiert die empirischen Ergebnisse unserer Studie und liefert eine strukturierte Analyse, wie sich die KI-Anwendungsfälle deutscher KI-Startups auf die Risikokategorien der EU KI-VO verteilen. Diese bilden die Grundlage für die Bewertung branchenspezifischer Auswirkungen und die Identifizierung von Bereichen, in denen gezielte Unterstützung und Kapazitätsaufbau am dringendsten benötigt werden.

Wir gliedern jede Beobachtung wie folgt:

- Wichtigste **Ergebnisse**
- Politische **Implikationen** unserer Erkenntnisse für die EU-Kommission, die deutsche Bundesregierung und die Bundesländer sowie für Startups und die Branche im weiteren Sinne
- **Empfehlungen** für diese Interessengruppen
- Systematische **Limitationen** der Studie und spezifische Limitationen der Ergebnisse.

Ergebnis 1: Der Anteil an KI-Startups mit hohem Risiko liegt deutlich über den ursprünglichen Annahmen der Europäischen Kommission

Abbildung 1 zeigt die Häufigkeit der einzelnen Risikoklassifizierungen in den Ergebnissen. Basierend auf dem KI-Anwendungsfall mit dem höchsten Risiko pro Startup werden 25,4 % (142 von 559) der analysierten Startups als Hochrisiko-KI eingestuft. Weitere 26,3 % (147 Startups) unterliegen Transparenzpflichten, während 48,3 % (270 Startups) in die Kategorie mit minimalem oder keinem Risiko fallen. Bemerkenswert ist, dass bei keinem der analysierten Startups festgestellt wurde, dass es verbotene KI-Systeme entwickelt.

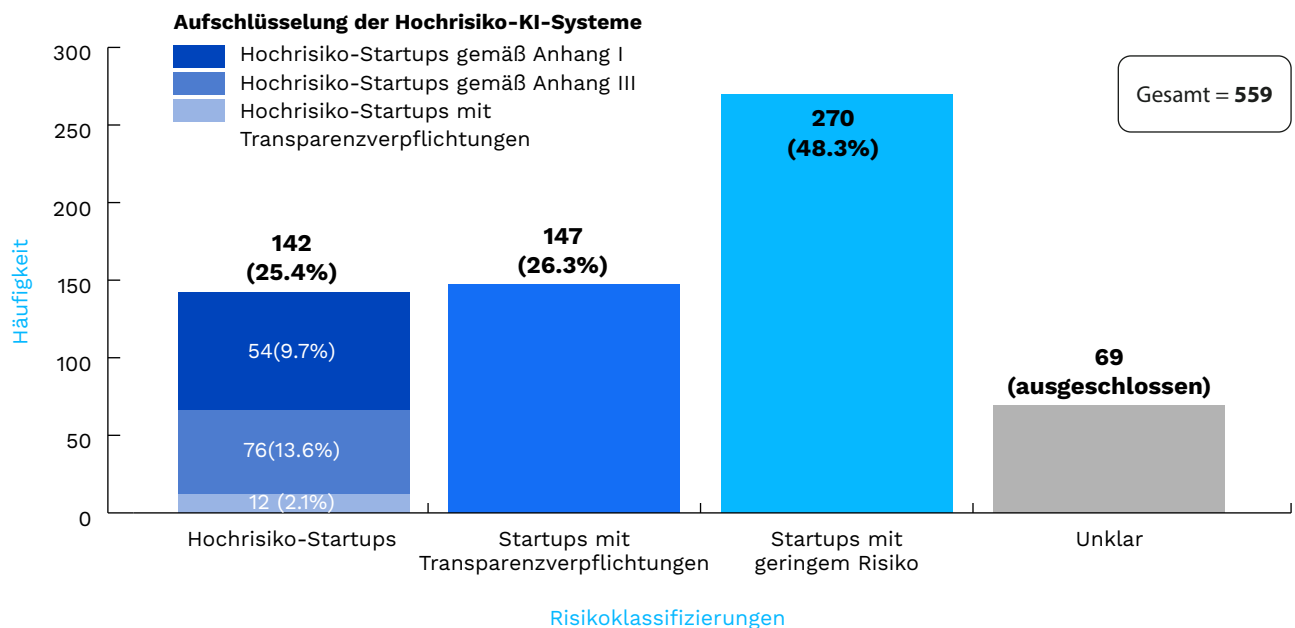


Abbildung 1: Verteilung der Risikoklassifizierung von KI-Startups

Implikationen

Akteure	Implikationen
EU-Kommission und die deutsche Bundesregierung	Die Folgenabschätzung der EU-Kommission aus dem Jahr 2021, in der geschätzt wurde, dass 5 % bis 15 % der Anwendungsfälle ein hohes Risiko darstellen würden, trifft auf deutsche KI-Startups nicht zu.
Bundesländer / Regionen	- Aufsichtsbehörden werden wahrscheinlich mit einer größeren Anzahl von Hochrisiko-Systemen konfrontiert sein als ursprünglich geplant. - Angesichts der umfangreichen Anforderungen des Entwurfs des Verhaltenskodex (CoP) für transparente KI-Systeme (z. B. Training, Tools, Dokumentation) werden die Aufsichtsbehörden zudem voraussichtlich eine größere Anzahl von KI-Systemen mit Transparenzpflichten überwachen müssen als erwartet (weitere Einzelheiten finden sich unter Ergebnis 3).

Empfehlungen

Akteure	Empfehlungen
EU / AI Office	• Das Impact Assessment der Europäischen Kommission von 2021 sollte überarbeitet werden, einschließlich der mit der KI-Verordnung verbundenen Compliance-Kosten. Mitgliedstaaten sollten ermutigt werden, ähnliche Studien in ihren Zuständigkeitsbereichen durchzuführen. • Falls erforderlich, sollten unterstützende Maßnahmen (einschließlich Finanzierungsmechanismen) angepasst werden, um der „tatsächlichen/ beobachteten“ regulatorischen Belastung Rechnung zu tragen.
Bundesregierung / Länder / Regionen	• Zuständige Aufsichtsbehörden müssen ausreichend ausgestattet sein, um das Volumen regulierter KI-Systeme bewältigen zu können (z. B. Bundesnetzagentur). • Es sollte sichergestellt werden, dass Startups mit Anwendungsfällen mit minimalen oder keinem Risiko nicht vernachlässigt werden. Sie sind nach wie vor erheblichen regulatorischen Belastungen durch andere Regelwerke (z. B. DSGVO, Haftung und geistiges Eigentum) ausgesetzt und sollten bei Maßnahmen zur Förderung von KI-Innovationen berücksichtigt werden.
Industrie / Startups	• Die Risikoklassifizierung der KI-Verordnung sollte frühzeitig in die Produktplanung integriert werden, um „kostspielige Überraschungen“ und die mögliche Notwendigkeit nachträglicher Überarbeitungen zu vermeiden. • Investoren sollten die Kosten für die Einhaltung regulatorischer Vorschriften in Bewertungsmodelle und Finanzierungsrunden einbeziehen, um sicherzustellen, dass deutsche Startups wettbewerbsfähig bleiben und keine Ressourcen vom Produkt auf die Compliance umleiten.

Limitationen

- **Interpretative limitation:** Given the absence of guidance, we did not consider whether any use cases might benefit from derogations in Article 6(3) of the AI Act.
- **Umfang der Analyse:** Unsere Ergebnisse beziehen sich auf deutsche Startups und lassen sich möglicherweise nicht auf andere Länder oder Unternehmen anderer Größe übertragen.
- **Informationsgrundlage:** Wir haben die Anwendungsfallbeschreibungen von Startup-Websites extrahiert und mithilfe von LLMs zusammengefasst. Diese Beschreibungen können unvollständig, zu Marketingzwecken vereinfacht oder nicht präzise genug sein, um den Einsatzkontext oder die regulatorischen Auswirkungen vollständig zu erfassen.
- **Methodik:** Diese Klassifizierung kombiniert automatisierte LLM-Analysen mit menschlicher Überprüfung. Dies ermöglicht zwar Skalierbarkeit, jedoch weisen beide Ansätze eine gewisse Fehlerquote auf.
- **Detaillierungsgrad der Überprüfung:** Die Analyse schlüsselt weder Unterkategorien innerhalb von Anhang I oder Anhang III der EU KI-VO auf, noch unterscheidet sie zwischen einzelnen Transparenzpflichten gemäß Artikel 50.
- **Interpretative Einschränkung:** Angesichts fehlender Leitlinien zum Zeitpunkt der Auswertung haben wir nicht geprüft, ob bestimmte Anwendungsfälle von Ausnahmeregelungen gemäß Artikel 6 Absatz 3 der KI-Verordnung profitieren könnten.

Ergebnis 2: Kontext und fehlende Details schränken die vollautomatische Risikoklassifizierung ein und erschweren die manuelle Überprüfung

Insgesamt 69 Startups konnten aufgrund unvollständiger Informationen zu ihren Anwendungsfällen (selbst nach manueller Bewertung aller dieser Anwendungsfälle) nicht abschließend klassifiziert werden. Diese Startups wurden von der weiteren Analyse ausgeschlossen. **Tabelle 1** enthält Beispiele für Anwendungsfälle, die wir nicht klassifizieren konnten, sowie die Gründe für unsere Unsicherheit.

Beispiele für Anwendungsfälle, bei denen Unsicherheiten hinsichtlich der Auslegung des Gesetzes oder fehlende Informationen über das KI-System eine Klassifizierung verhinderten

Beispiel 1

Batterieanalyse

Verwendungszweck:

Eine Plattform zur prädiktiven Analyse von Batterien, die darauf ausgelegt ist, die Sicherheit, Leistung und Lebensdauer von Lithium-Ionen-Batteriesystemen zu verbessern.

Unsicherheit:

Die Risikoklassifizierung hängt stark vom Anwendungskontext ab. Erfüllt ein solches System, eine Sicherheitsfunktion, könnte es als Hochrisiko-System eingestuft werden, sei es durch den Einsatz in kritischer Infrastruktur oder die Verwendung in Produkten, die unter die Regelungen von Anhang I fallen.

Beispiel 2

Datenaggregation

Verwendungszweck:

Ein KI-Agent, der kontinuierlich globale Nachrichtenquellen scannt und analysiert, um Echtzeit-Informationen zu aktuellen Ereignissen und Entwicklungen zu liefern.

Unsicherheit:

Aus der Beschreibung des Anwendungsfalls geht nicht eindeutig hervor, ob das System synthetische Inhalte generiert (was eine Transparenzpflicht nach sich ziehen würde) oder lediglich relevante Nachrichteninhalte identifiziert.

Beispiel 3

Bewegungserfassung

Verwendungszweck:

Das System bietet hochpräzise Echtzeit-Tracking-Systeme und Bewegungsanalysen von Menschen auf der Grundlage einfacher Videoaufzeichnungen und wird von verschiedenen Betreibern in verschiedenen Bereichen eingesetzt, darunter Sportorganisationen und medizinische Kliniken.

Unsicherheit:

Es ist unklar, ob das System auch in der Lage ist, Zielpersonen zu identifizieren.

Tabelle 1: Beispiele für unklare Klassifizierungen

Implikationen

Akteure	Implikationen
EU / AI Office	- Die Annahme der KI-Verordnung, dass jeder Anwendungsfall einen „beabsichtigten Verwendungszweck“ hat, trifft möglicherweise nicht immer zu. Die Risikoklassifizierung kann sich je nach Einsatzkontext und Branche verschieben. Tatsächlich zeigt unsere Studie, dass Startups häufig branchenübergreifend tätig sind (siehe Ergebnis 4 für Details). - Eine unklare Risikoklassifizierung kann verschiedene Ursachen haben, darunter KI-Systeme mit mehreren Verwendungszwecken, unvollständige oder missverständliche Beschreibungen von Anwendungsfällen oder ungenaue Auslegungen der KI-Verordnung.
Bundesregierung / Bundesländer / Regionen	Ohne klare Leitlinien könnte eine dezentrale Aufsicht zu widersprüchlichen Auslegungen der KI-Verordnung und voneinander abweichenden Einstufungen durch die verschiedenen Aufsichtsbehörden führen.
Industrie	Dasselbe KI-System kann je nach Anwendungsbereich oder Unternehmensfunktion in unterschiedliche Risikoklassen fallen (z. B. kann die Dokumentenanalyse im Personalwesen als hohes Risiko gelten, im Marketing hingegen als mit minimalem oder keinem Risiko).

Empfehlungen

Akteure	Empfehlungen
EU / AI Office	Es sollten klare Leitlinien dazu bereitgestellt werden, wie der „beabsichtigte Verwendungszweck“ grundsätzlich auszulegen ist, einschließlich für Mehrzweck- oder modulare KI-Systeme sowie Multi-Agenten-Systeme.
Bundesregierung / Länder / Regionen	- Bei der Marktüberwachung sollte ein „Guidance-First“-Ansatz verfolgt werden, der gewissenhafte Bemühungen honoriert und die technischen Herausforderungen bei der Risikoklassifizierung berücksichtigt. - Koordinierungsmechanismen sollten eingerichtet werden, um eine einheitliche Auslegung der Risikoklassifizierung bei allen Aufsichtsbehörden sicherzustellen. - Lokale Initiativen sollten unterstützt werden, die zur Klärung der Risikoklasse beitragen und Rechtsunsicherheit für die Industrie und den öffentlichen Sektor ausräumen können.
Industrie / Startups	- Betreiber und Anbieter sollten den spezifischen Zweck und den Nutzungskontext jedes KI-Systems berücksichtigen, da dies Konsequenzen in Form zusätzlicher/anderer Verpflichtungen gemäß der KI-Verordnung nach sich ziehen kann⁵. - Produktteams in Startups sollten Änderungen am Produkt oder am Anwendungsfall bewusst steuern, da Iterationen die Risikoklassifizierung verändern können.

Limitationen

- **Informationsgrundlage:** Wir haben die Anwendungsfallbeschreibungen von den Websites aller Startups gesammelt und mithilfe von LLMs zusammengefasst. Diese Beschreibungen können unvollständig, zu Marketingzwecken vereinfacht oder nicht präzise genug sein, um den Einsatzkontext oder die regulatorischen Auswirkungen vollständig zu erfassen.
- **Interpretationsunsicherheit:** Die rechtliche Auslegung mehrerer Konzepte der KI-Verordnung ist noch ungeklärt und kann sich durch Leitlinien oder Rechtsprechung weiterentwickeln. Bis dahin besteht die Möglichkeit, dass menschliche Prüfer das Gesetz falsch verstanden haben.

Infobox

Die entscheidende Rolle der menschlichen Überprüfung bei der Risikoklassifizierung

Die genaue Einstufung von KI-Systemen gemäß der EU KI-VO ist kein rein schematischer Prozess. Zwar stehen zunehmend Leitfäden und Auslegungshilfen zur Verfügung, doch unsere Analyse zeigt, dass eine zuverlässige Einstufung die kombinierte Expertise von Juristinnen und Juristen als auch von Technologieexpertinnen und -experten erfordert. Die Risikoklasse eines KI-Systems hängt häufig von feinen Details seines tatsächlichen Einsatzkontexts und den technischen Details des Produkts ab, die durch allgemeine Leitlinien allein nicht erfasst werden können.

In der Praxis stießen wir wiederholt auf Grenzfälle, in denen identische technische Funktionen je nach Anwendungsbereich oder letztendlicher Nutzung in unterschiedliche Kategorien fallen konnten. So können beispielsweise Lösungen zur vorausschauenden Wartung für Fertigungsmaschinen als mit minimalen oder keinen Risiken eingestuft werden, während derselbe Ansatz bei kritischen Infrastrukturen wie Windkraftanlagen eine Risikoklassifizierung als hohes Risiko nach sich ziehen kann. Ebenso führen KI-Tools im Personalwesen häufig zu Unsicherheit. So könnten KI-Lösungen für die Schichtplanung oder das Personalmanagement ein minimales oder kein Risiko darstellen, wenn sie auf mitarbeiterunabhängigen Parametern wie dem prognostizierten Kundenaufkommen basieren, aber leicht zu einem Hochrisiko-System werden, wenn sie auf Mitarbeitermerkmalen beruhen.

Diese Herausforderungen werden durch die Verweise der KI-Verordnung auf andere EU-Rechtsrahmen noch verstärkt. Anhang I beispielsweise integriert Verpflichtungen aus mehreren branchenspezifischen Verordnungen, was zu Grauzonen in der Auslegung führt. Die Robotik veranschaulicht diese Komplexität: Nicht alle KI-gestützten oder KI-fähigen Maschinen müssen gemäß der aktuellen Maschinenrichtlinie einer Bewertung durch Dritte unterzogen werden; einige müssen dies gemäß der künftigen Maschinenverordnung möglicherweise; während andere wiederum anderen branchenspezifischen Gesetzen unterliegen könnten, wie beispielsweise denen für Fahrzeuge oder landwirtschaftliche Geräte.

Diese Beispiele verdeutlichen die inhärenten Grenzen einer Recherche, die sich ausschließlich auf öffentlich zugängliche Informationen stützt. In unserer Studie konnten eine Reihe von Anwendungsfällen aufgrund fehlender Details zur Einsatzumgebung oder ungeklärter rechtlicher Auslegungen nicht abschließend klassifiziert werden. Dies unterstreicht die entscheidende Rolle des Menschen im Prozess der Risikoklassifizierung und die Grenzen automatisierter Ansätze oder zu stark vereinfachten Checklisten. Menschliches Urteilsvermögen – fundiert sowohl in rechtlicher Auslegung als auch in technischem Fachwissen – bleibt vorerst unverzichtbar, um eine genaue und kontextbezogene Anwendung der EU KI-VO zu gewährleisten.

Ergebnis 3: Die Verbreitung von GPAI-Modellen macht die Transparenzpflichten zur vorherrschenden regulatorischen Grundlage für KI-Startups

Unsere Studie zur deutschen Startup-Landschaft 2025 hat gezeigt, dass fast jedes dritte deutsche Startup generative KI-Lösungen entwickelt, wobei die Zahl dieser Startups im Vergleich zu 2024 um 130 % gestiegen ist. Erwartungsgemäß unterliegen voraussichtlich 59 % der Generative KI-Startups in unserer Studie – die auf der „German Startup Landscape Study 2024“ basiert – Transparenzpflichten. Darüber hinaus wurden 12 % dieser Startups zudem als Hochrisiko eingestuft. **Abbildung 2** zeigt die Verteilung von insgesamt 111 generativen KI-Startups auf die einzelnen Risikokategorien.

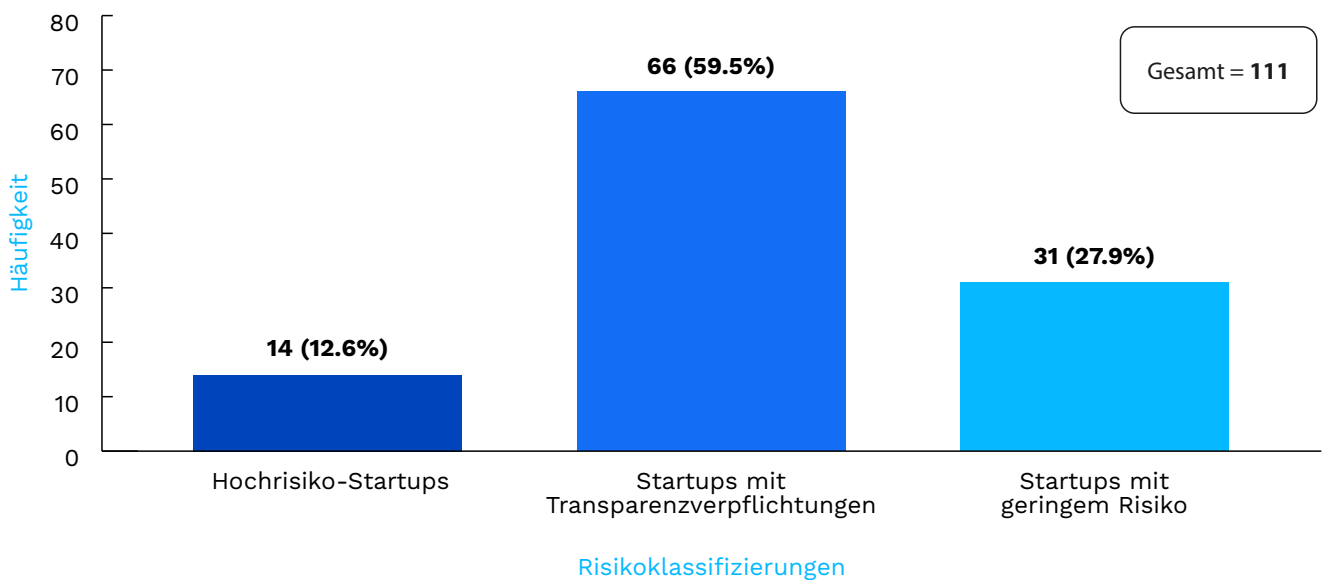


Abbildung 2: Risikoklassifizierung nach Startups im Bereich Generative KI

Implikationen

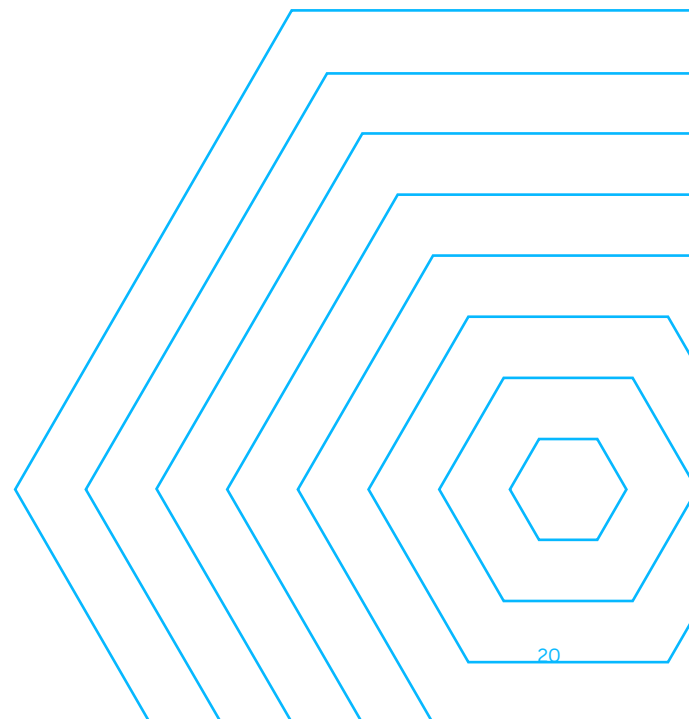
Akteure	Implikationen
EU / AI Office	• Transparenzpflichten dürften einen größeren (und potenziell wachsenden) Anteil deutscher Startups betreffen, da die Nutzung von LLM im Laufe der Zeit zunimmt. • Inwieweit Startups in der Lage sind, die KI-Verordnung und insbesondere die Anforderungen für Hochrisiko-KI-Systeme einzuhalten, wird maßgeblich davon abhängen, inwieweit Anbieter von KI-Modellen für allgemeine Zwecke (GPAI) die Verpflichtungen zur Weitergabe von Informationen und Ressourcen gemäß dem Verhaltenskodex für KI-Modelle für allgemeine Zwecke einhalten.
Industrie / Startups	• Die Leistungsfähigkeit und die Einhaltung gesetzlicher Vorschriften deutscher KI-Startups hängen zunehmend von Anbietern von KI-Modellen für allgemeine Zwecke außerhalb der EU ab. KI-Startups müssen die potenziellen Abhängigkeiten abwägen, die sie eingehen, wenn sie ein GPAI-Modell als Grundlage für ihr KI-System oder ihr Geschäftsmodell wählen.

Empfehlungen

Akteure	Empfehlungen
EU / AI Office	• Die Annahme des Omnibus-Vorschlags zur Aufschiebung der Transparenzpflichten sollte erwogen werden, da der Verhaltenskodex für transparente KI-Systeme erst im Sommer 2026 fertiggestellt sein wird. • Die Anwendung des Verhaltenskodexes für KI-Modelle für allgemeine Zwecke sollte beobachtet werden und sichergestellt sein, dass nachgelagerte Anbieter ausreichende Informationen erhalten, um die Einhaltung der Vorschriften zu gewährleisten. • Die Entwicklung von in der EU ansässigen Anbietern von KI-Modellen für allgemeine Zwecke sollte unterstützt werden, um eine durchgängige Einhaltung der Vorschriften entlang der gesamten KI-Wertschöpfungskette sicherzustellen.
Industrie / Startups	• Es sollte sichergestellt werden, dass die für die nachgelagerte Einhaltung erforderlichen technischen Details vom Anbieter von KI-Modellen für allgemeine Zwecke als Teil der Modellauswahlkriterien bereitgestellt werden, insbesondere wenn der Anwendungsfall mit einem hohem Risiko verbunden ist.

Limitationen

- **Definition eines „Generativen KI“-Startups:** Wir stützen uns auf den in unserer deutschen Landschaftsstudie 2024 verwendeten Klassifizierungsprozess, um zu bestimmen, wann ein Startup als Generatives KI-Startup gilt und haben keine eigene Methode entwickelt.



Ergebnis 4: Die Compliance-Belastung betrifft bestimmte Branchen und Unternehmensfunktionen überproportional

Die Compliance-Belastung für deutsche Startups ist nicht gleichmäßig verteilt. Hochrisiko-KI-Systeme und Transparenzpflichten konzentrieren sich häufig auf bestimmte Branchen und Unternehmensfunktionen.

Abbildung 3 zeigt die Häufigkeit aller Risikokategorien in den einzelnen Branchen. Die Abbildung zeigt Daten für 369 Startups, denen wir manuell eine Branchenzugehörigkeit zuordnen konnten.

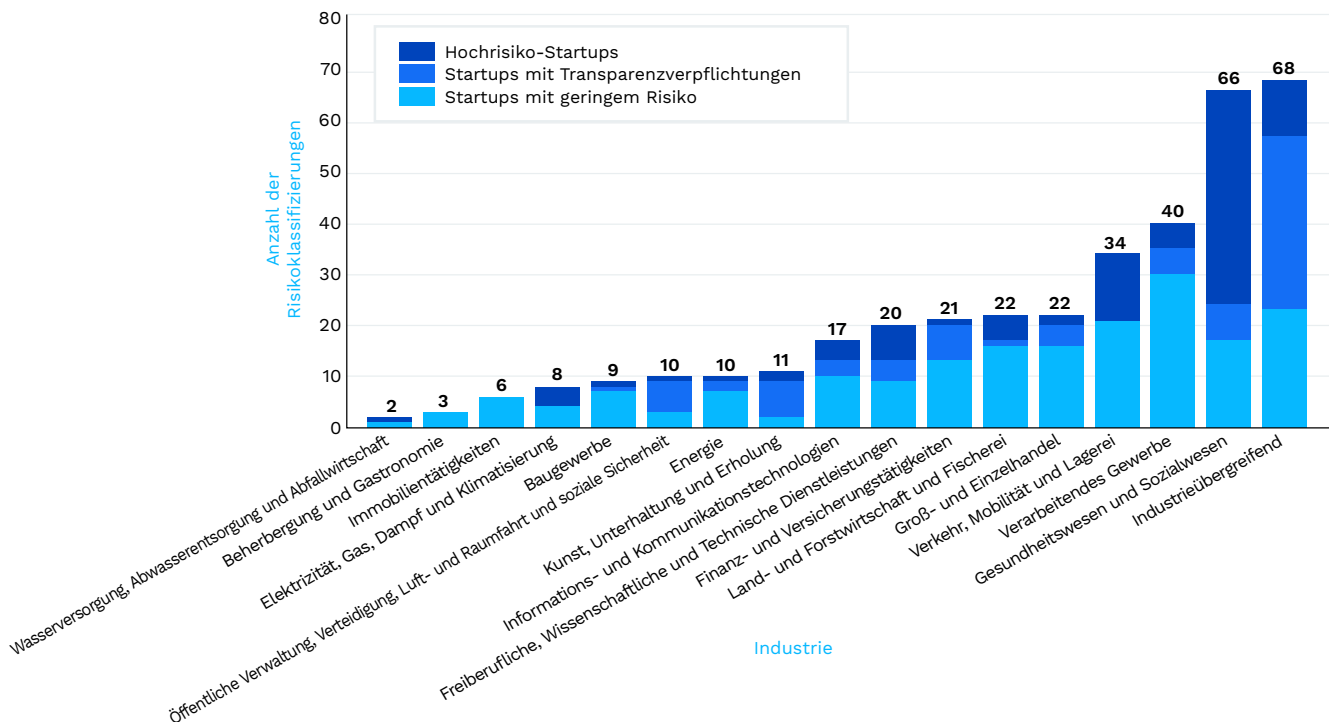


Abbildung 3: Verteilung der Risikoklassen nach Branchen

Wir beobachten eine Konzentration von Startups mit hohem Risiko in Branchen wie der Gesundheitsbranche (z. B. KI in Medizinprodukten), dem Transportwesen (z. B. autonomes Fahren) und in „branchenübergreifenden“ Anwendungsfällen (typischerweise Startups, deren Produkt branchenunabhängig ist, wie z. B. im Personalwesen). Ebenso stellen wir fest, dass ein erheblicher Anteil der branchenübergreifend tätigen Startups voraussichtlich Transparenzpflichten unterliegt.

Abbildung 4 zeigt die Verteilung der Risikoklassifizierungen nach Unternehmensfunktion. Die Grafik zeigt Daten für 282 Anwendungsfälle, denen wir manuell eine Unternehmensfunktion zuweisen konnten.

Die höchste Konzentration an Startups mit hohem Risiko findet sich in Unternehmensfunktionen wie dem Personalwesen, was einen wachsenden Trend widerspiegelt, KI-Systeme einzusetzen, um die Rekrutierungskapazitäten von Unternehmen auszubauen. Ebenso lässt sich erkennen, dass sich Transparenzpflichten auf Unternehmensfunktionen wie Marketing und Kundensupport konzentrieren, wo LLM-basierte KI-Systeme zunehmend von Unternehmen genutzt werden. Zu den Unternehmensfunktionen mit einem hohen Anteil an KI-Systemen mit minimalem oder keinem Risiko gehören Produktion, Forschung und Entwicklung sowie Betrieb, was auf Bereiche hinweist, für die die KI-Verordnung keine zusätzlichen Anforderungen stellt.

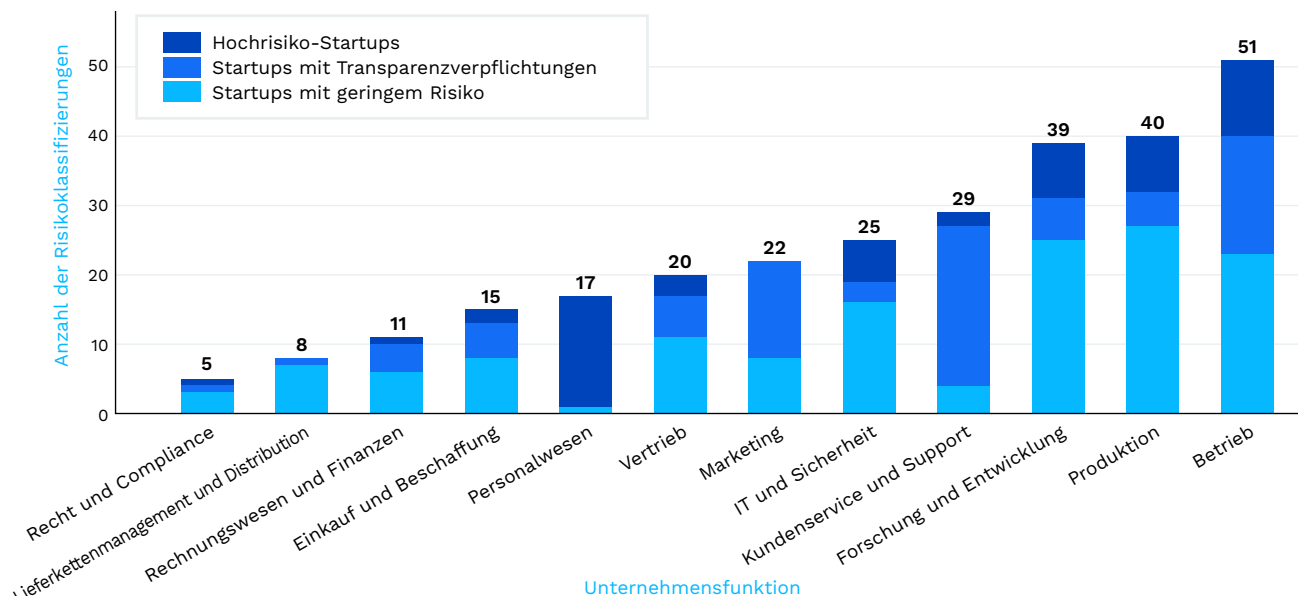


Abbildung 4: Verteilung der Risikoklassen nach Unternehmensfunktion

Implikationen

Akteure	Implikationen
EU / AI Office	Einheitliche und „horizontale“ (d. h. branchenunabhängige) Leitlinien zur Compliance spiegeln möglicherweise nicht die Realität wider, weil sich die regulatorischen Belastungen auf bestimmte Branchen und Unternehmensfunktionen konzentrieren.
Länder / Regionen	Einige Branchen und Bereiche werden unverhältnismäßig mehr regulatorische Unterstützung und Aufsicht durch Aufsichtsbehörden und benannte Stellen benötigen.
Industrie / Startups	Einige Branchen und Unternehmensbereiche fallen strukturell bedingt eher in die Hochrisiko-Kategorie. So erfordern beispielsweise Anwendungsfälle im Personalwesen und im operativen Betrieb eine stärkere Aufsicht.

Empfehlungen

Akteure	Empfehlungen
EU / AI Office	Eine Entwicklung branchenspezifischer Leitlinien an der Schnittstelle zwischen der KI-Verordnung und anderen Vorschriften, wie z. B. für Medizinprodukte, sollte gefördert werden.
Mitgliedstaaten / Regionen	- Die Unterstützung bei der Einhaltung der Vorschriften muss auf die jeweilige Branche und die Unternehmensfunktion zugeschnitten sein. So könnten beispielsweise viele Unternehmen von Leitlinien zum Thema „KI-Verordnung im Personalwesen“ profitieren. - Es sollten Untersuchungen zu KI-zentrierten KMUs in den Bundesländern durchgeführt werden, um ein umfassenderes Bild der branchenspezifischen Risikoprofile zu erhalten.
Industrie / Startups	- Der Weiterbildung und Umschulung von Mitarbeitern, die KI-Systeme entwickeln und einsetzen, sollte Priorität eingeräumt werden. - Branchenverbände und KMU-Zentren auf Landesebene sollten zusammenarbeiten, um die Kosten für die Rechtsauslegung und die technische Compliance zu senken.

Limitationen

- **Umfang der Analyse:** Unsere Ergebnisse beziehen sich auf deutsche Startups und lassen sich möglicherweise nicht auf andere Rechtsordnungen oder Unternehmensgrößen übertragen.

Ergebnis 5: Deutsche Bundesländer weisen ähnliche Profile hinsichtlich der Risikoklasse auf, unterscheiden sich jedoch in ihrer Verteilung auf die einzelnen Branchen

In den deutschen Bundesländern mit bedeutender Startup-Aktivität ist die Verteilung der Risikoklassifizierungen nach der KI-Verordnung weitgehend ähnlich, trotz großer Unterschiede in der absoluten Anzahl der Startups. Die Zusammensetzung der Startup-Aktivitäten über die verschiedenen Branchen hinweg unterscheidet sich jedoch deutlich zwischen den Bundesländern, was zu unterschiedlich starken Auswirkungen der KI-VO führt.

Abbildung 5 zeigt die Verteilung der Risikoklassifizierung nach Bundesländern und **Abbildung 6** konzentriert sich auf die Verteilung der fünf Bundesländer mit der höchsten Gesamtzahl an KI-Systemen.

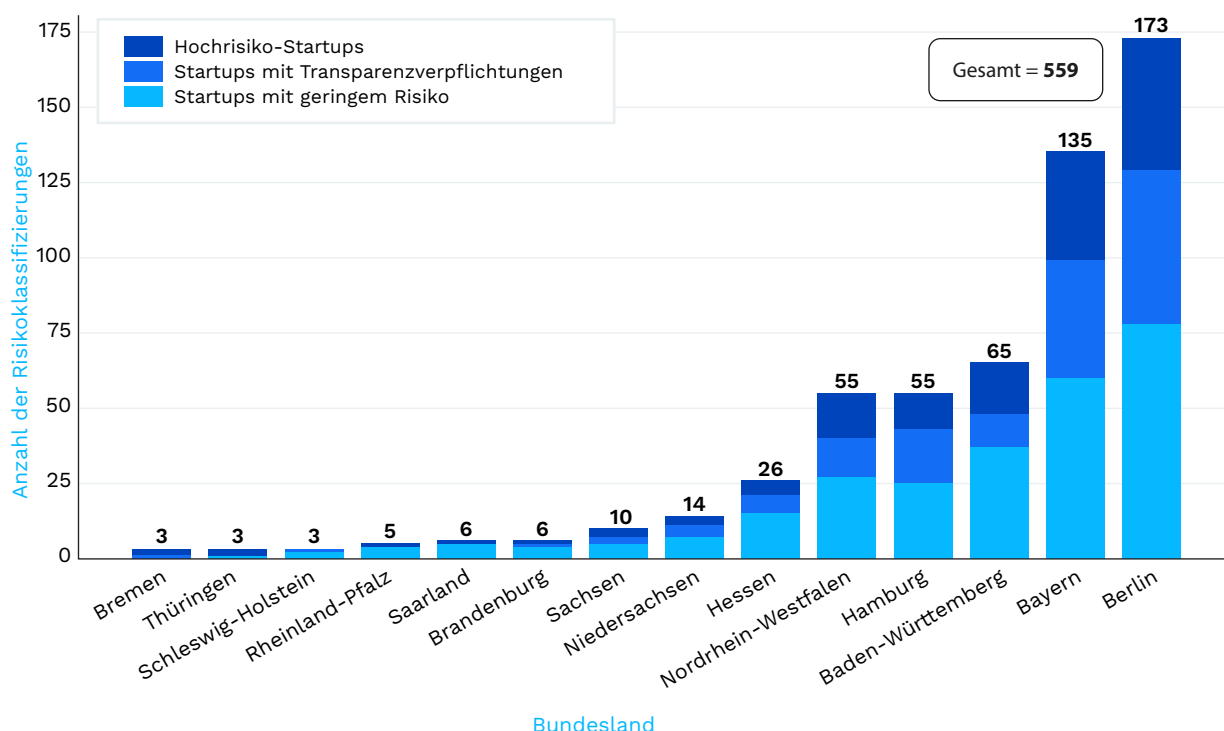


Abbildung 5: Verteilung der Risikoklassifizierung nach Bundesländern

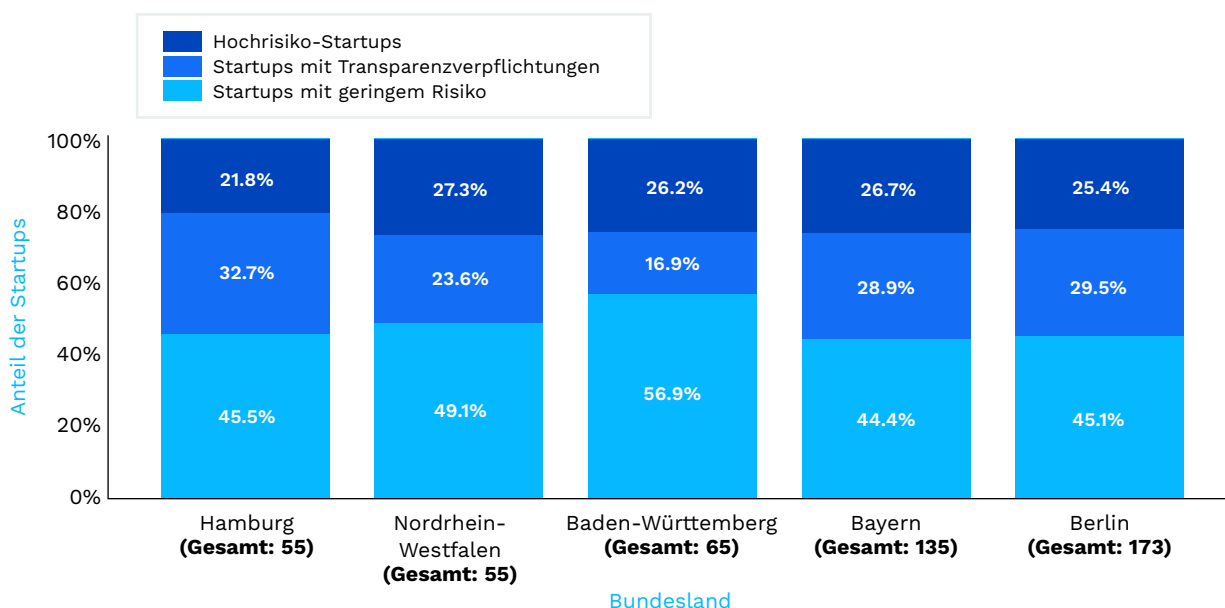


Abbildung 6: Verteilung der Risikoklassifizierung nach Bundesländern mit der höchsten Startup-Aktivität

Unter den fünf Bundesländern mit der höchsten Anzahl an KI-Startups ist die Verteilung von Startups mit minimalem oder keinem Risiko, solchen mit Transparenzpflichten und solchen mit hohem Risiko weitgehend ähnlich. In den übrigen Bundesländern ist die Gesamtaktivität der Startups zu gering, um aussagekräftige Vergleiche anzustellen.

Im Vergleich dazu sieht die sektorale Verteilung sehr anders aus. **Abbildung 7** zeigt eine Heatmap der 351⁶ KI-Startups über alle Bundesländer hinweg, für die mindestens acht Anwendungsfälle und Branchen vorliegen und für die eine manuelle Klassifizierung verfügbar sind.

Die Zahl in jeder Zelle gibt die Anzahl der Startups in der auf der x-Achse bezeichneten Branche an, die sich in dem auf der y-Achse bezeichneten Bundesland befinden. Die Farbskala der Heatmap reicht von Hell- bis Dunkelblau, wobei Hellblau für eine geringe Anzahl von Startups und Dunkelblau für eine hohe Anzahl steht. Zellen ohne Startups bleiben leer.

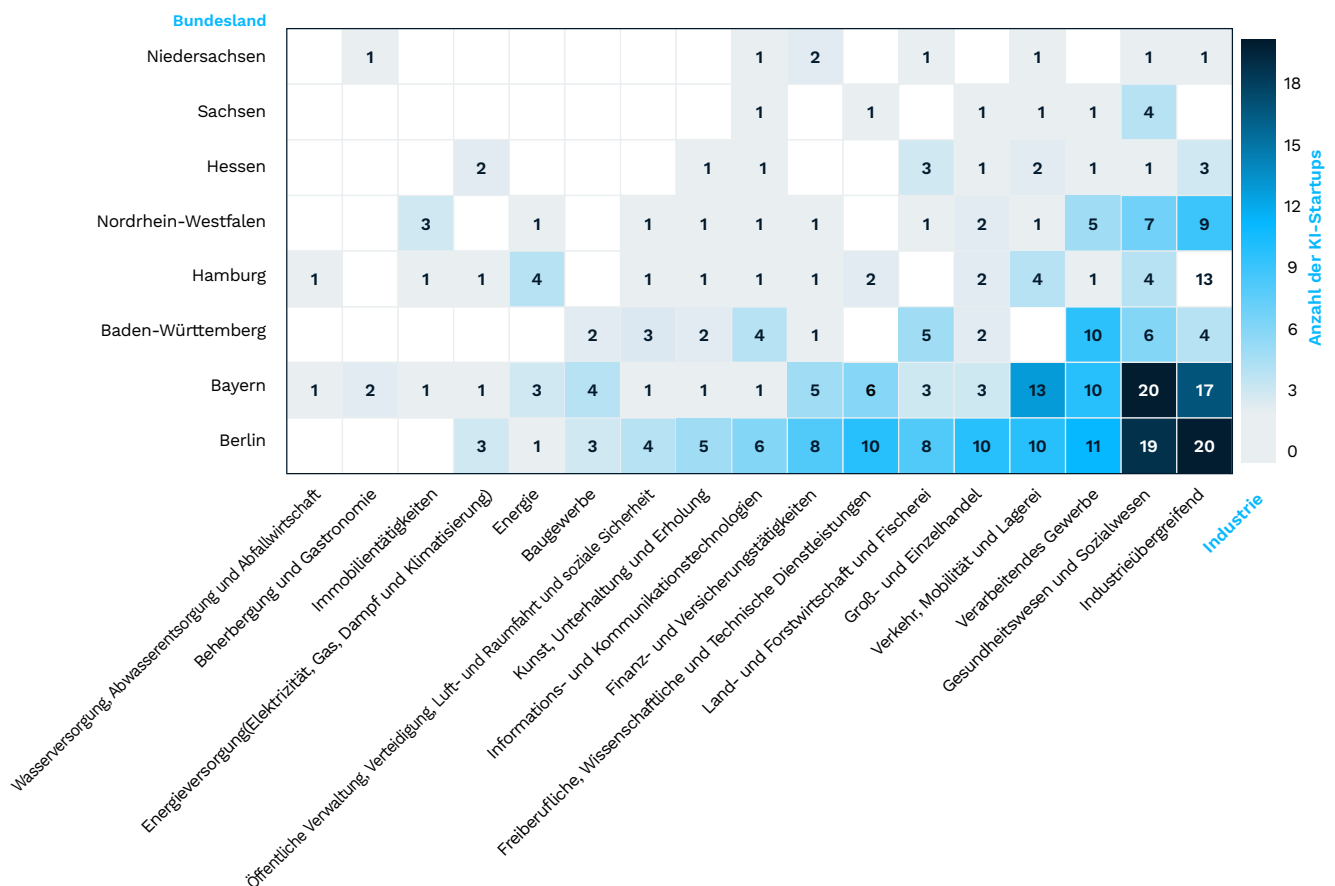


Abbildung 7: Heatmap der KI-Startup-Aktivitäten nach Bundesland und Branche

Die Bundesländer am unteren Ende der y-Achse (z. B. Bayern und Berlin) und die Branchen am rechten Ende der x-Achse (z. B. Gesundheitswesen und Sozialwesen) weisen eine höhere Anzahl an Startups auf. Bayern ist das einzige Bundesland, das in allen Branchen vertreten ist, während Berlin die höchste Gesamtzahl an Startups verzeichnet. Beide Bundesländer zeigen eine hohe Konzentration in den Bereichen Industrieübergreifend, Gesundheitswesen und Sozialwesen, Verarbeitendes Gewerbe sowie Verkehr, Mobilität und Lagerei. Andere Bundesländer weisen deutlich weniger diversifizierte Branchenprofile auf, die sich häufig auf das verarbeitende Gewerbe, den Energiesektor, den Immobiliensektor oder andere Branchen konzentrieren.

6 Acht Bundesländer wiesen weniger als acht Startups auf. Diese haben wir aus der vorliegenden Analyse ausgeschlossen. Für die übrigen Bundesländer haben wir nur Anwendungsfälle berücksichtigt, für die Branchenangaben verfügbar waren (siehe auch Abb. 3).

Implikationen

Akteure	Implikationen
Bundesländer / Regionen	Die regulatorischen Anforderungen unterscheiden sich erheblich zwischen den Bundesländern, was differenzierte Aufsichts-, Beratungs- und Unterstützungskapazitäten erfordert.
Industrie/ Startups	Viele Startups sind „industriübergreifend“ tätig, d. h., sie entwickeln branchenunabhängige Produkte und Dienstleistungen.

Empfehlungen

Akteure	Empfehlungen
Bundesländer / Regionen	- Einrichtung von Mechanismen zur länderübergreifenden Zusammenarbeit, die sich auf branchenspezifische KI-Anwendungsfälle konzentrieren und den Austausch von Fachwissen, Aufsichtskapazitäten und bewährten Verfahren ermöglichen. Ein Beispiel hierfür wäre eine Zusammenarbeit im verarbeitenden Gewerbe zwischen Berlin, Bayern und Baden-Württemberg - Beauftragung weiterer empirischer Analysen zur Branchenstruktur auf Landesebene, um die Aufsichtskapazitäten und Unterstützungsmaßnahmen besser auf die regionalen Risikoprofile abzustimmen. - Es sollte sich darauf eingestellt werden, dass der Regulierungsbedarf in Deutschland möglicherweise nicht ungleich verteilt ist und dass demnach Aufsichts- und Marktüberwachungsbehörden unterschiedlich ausgestattet sein müssen.
Industrie/ Startups	Stärkung der strukturierten Unterstützung für KI-Startups durch Branchenverbände und private Kapitalgeber, insbesondere in Hochrisiko-Bereichen oder branchenübergreifenden Bereichen, um den Zugang zu Compliance-Fachwissen und gemeinsamen Governance-Ressourcen zu verbessern. Diese Unterstützung könnte durch Verbände wie beispielsweise den VDMA, den ZVEI, den BDI oder den BVME angeboten werden.

Limitationen

- **Geringe Stichprobengröße:** In vielen Bundesländern ist die Gesamtzahl der Startups in unserer Studie gering, was aussagekräftige Vergleiche einschränkt.
- **Teilweise Abdeckung des Ökosystems:** Die Analyse umfasst ausschließlich KI-Startups. Daher geben die Ergebnisse nicht die gesamte regionale KI-Landschaft wieder, z. B. die Verbreitung von KI in KMU und Großunternehmen.

Ergebnis 6: Compliance-Kosten spiegeln sich (noch) nicht in ihrer Finanzierung wider

Trotz eines erheblichen Anteils an Startups, die als Hochrisiko eingestuft werden oder einer Transparenzpflicht unterliegen, spiegeln die Finanzierungsmuster mögliche Compliance-Kosten im Zusammenhang mit der KI-Verordnung noch nicht wider.

Abbildung 8 zeigt die durchschnittliche und die mediane Finanzierung⁷ je nach Risikoklassifizierung. Die durchschnittliche Finanzierung ist bei Startups mit Transparenzpflichten am höchsten und bei Startups mit minimalem oder keinem Risiko am niedrigsten. Die im Vergleich zu den Mittelwerten über alle Gruppen hinweg deutlich niedrigeren Mediane weisen auf eine stark asymmetrische Verteilung hin. Mit anderen Worten: Einige wenige Ausreißer-Startups erhalten einen deutlich größeren Anteil an der Finanzierung als die verbleibende Mehrheit.

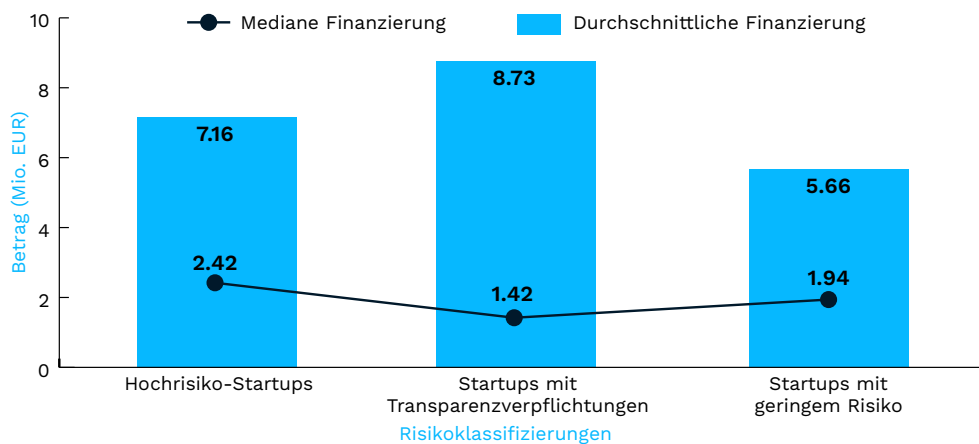


Abbildung 8: Durchschnittliche und mediane Finanzierung von Startups nach Risikoklassifizierung

Abbildung 9 zeigt die Gesamtfinanzierung, die alle Startups in einem bestimmten Bundesland erhalten haben. Die fünf Bundesländer mit der höchsten Gesamtfinanzierung sind dieselben Bundesländer mit der höchsten Anzahl an KI-Startups, wie in Abbildung 5 dargestellt, und zwar in derselben Reihenfolge, d. h. Hamburg, Nordrhein-Westfalen, Baden-Württemberg, Bayern und Berlin. **Abbildung 10** zeigt die Verteilung der Gesamtfinanzierung über die Risikokategorien der fünf Bundesländer mit der höchsten Startup-Aktivität.

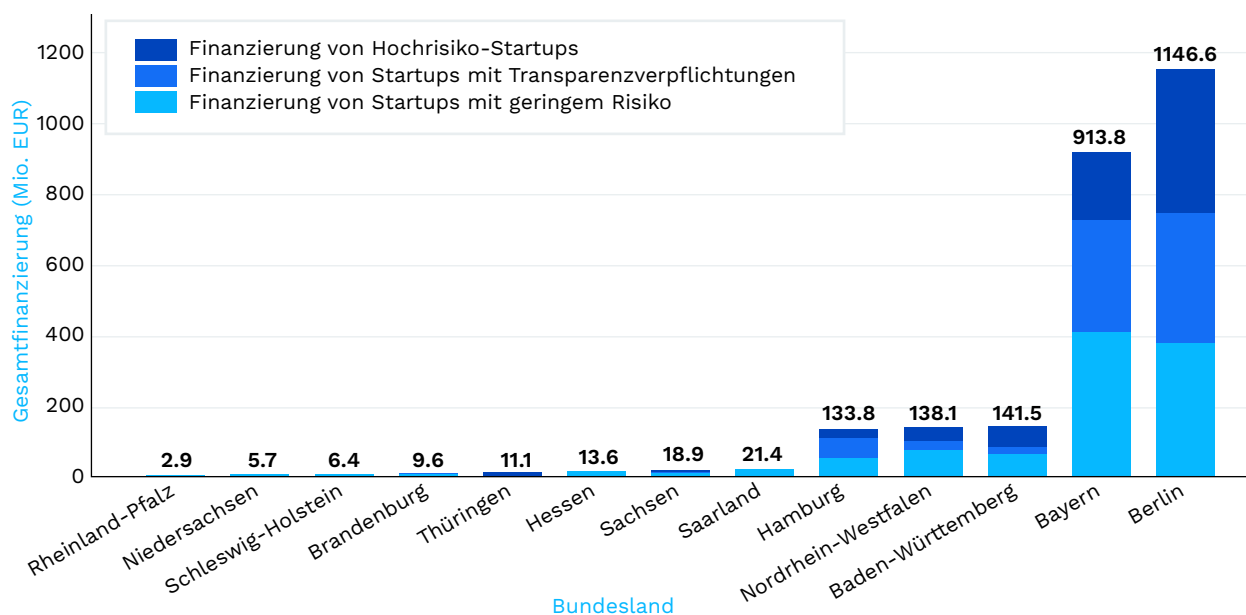


Abbildung 9: Gesamtfinanzierung nach Bundesland und Risikoklassifizierung

⁷ Die Daten stammen von dealroom.co und wurden im Juli 2024 durch interne Datenquellen ergänzt. Sie geben das insgesamt eingeworbene Kapital wieder, das in erster Linie von Investoren aus dem privaten Markt stammt.

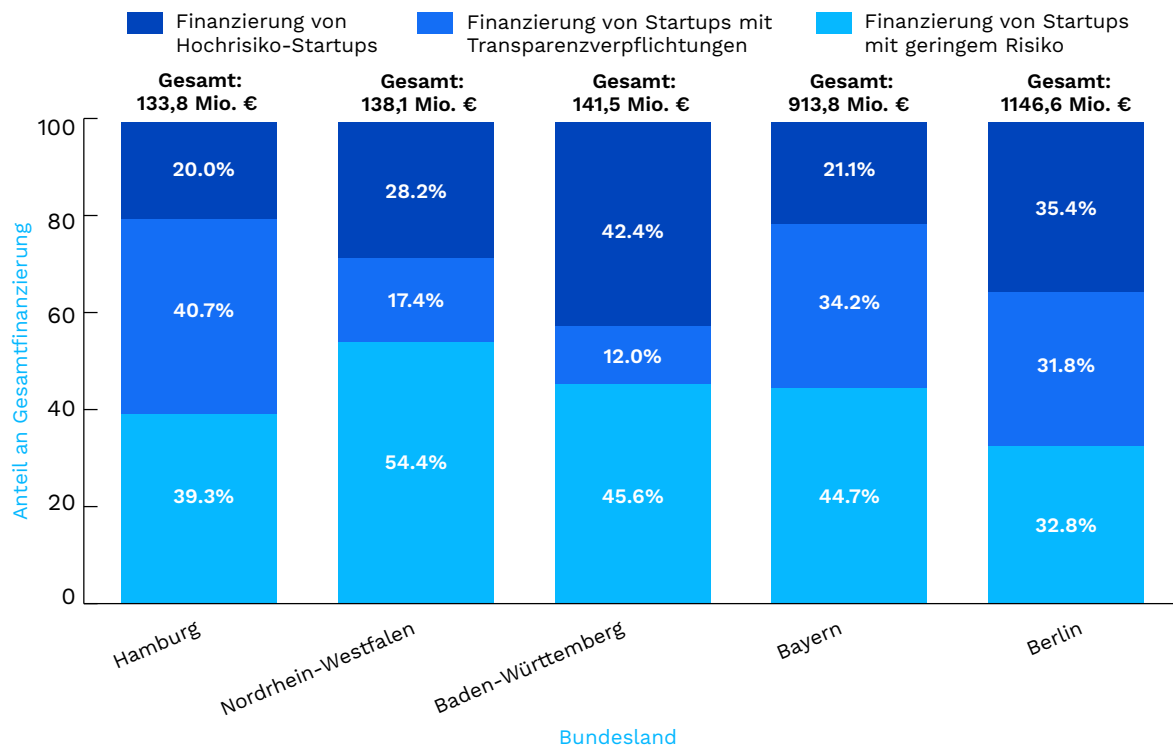


Abbildung 10: Verteilung der Fördermittel nach Risikoklassifizierung für die Bundesländer mit der höchsten Startup-Aktivität

Angesichts der asymmetrischen Verteilung der Finanzierung (d. h. ein kleiner Prozentsatz der Startups erhält den größten Anteil der Finanzierung) und der vergleichsweise geringen Anzahl von KI-Startups aus Bundesländern außer Berlin und Bayern sind wir der Ansicht, dass diese Daten mit Vorsicht interpretiert werden sollten und nicht als Hinweis auf strukturelle Muster bei den Investitionen zu verstehen sind (siehe Limitationen unten).

Implikationen

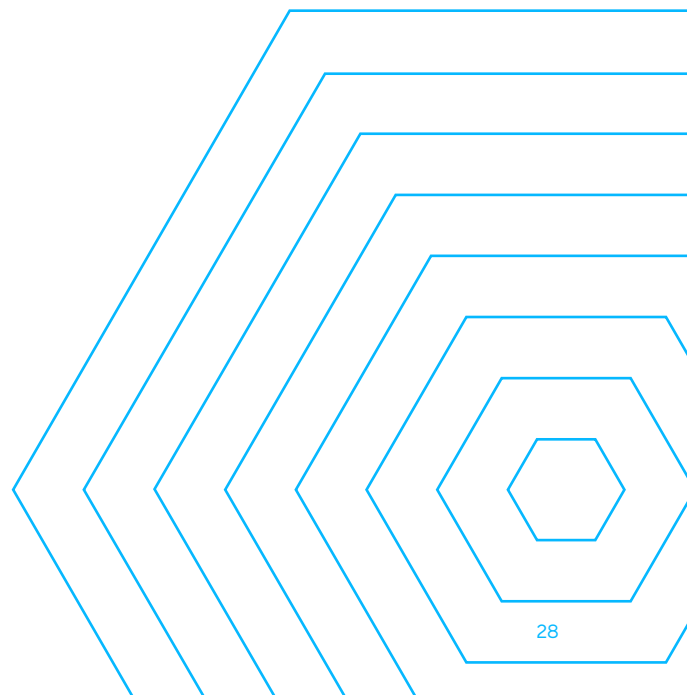
Akteure	Implikationen
Bundesländer / Regionen	Setzt man die moderate Median-Finanzierung pro Startup über alle Risikoklassen hinweg ins Verhältnis zu den erwarteten Compliance-Kosten von rund 200.000 bis 300.000 Euro für ein Hochrisiko-KI-System (Renda <i>et al.</i>, 2021, S. 152–153), lässt sich die relative regulatorische Belastung für Startups abschätzen.
Branche / Startups	- Die stark ungleich verteilte Finanzierung in allen Klassen deutet darauf hin, dass gut finanzierte Startups die Compliance-Anforderungen möglicherweise besser bewältigen können als die vielen Startups mit moderater Finanzierung. - Die Toleranz von Risikokapitalgebern gegenüber Compliance-Kosten dürfte sich je nach Sektor unterscheiden (z. B. Biotechnologie vs. EdTech), was die Chancen von KI-Startups auf die Beschaffung von Finanzmitteln und damit ihr Überleben beeinflussen könnte.

Empfehlungen

Akteure	Empfehlungen
Mitgliedstaaten / Regionen	• Damit das Überleben von Startups weniger davon abhängt, ob die Compliance-Kosten für sie tragbar sind, sollten Landesregierungen (mit hohem Regulierungsaufwand) umfassende Unterstützung anbieten. • Weitere Untersuchungen sollten beauftragt werden, um zu ermitteln, ob sich die Risikoklassifizierung auf die Finanzierungsanforderungen auswirkt und in welchem Umfang genau.
Industrie / Startups	• Es sind weitere Studien erforderlich, um die Bereitschaft von Risikokapitalgebern zur Finanzierung von Hochrisiko-KI-Startups zu verstehen.

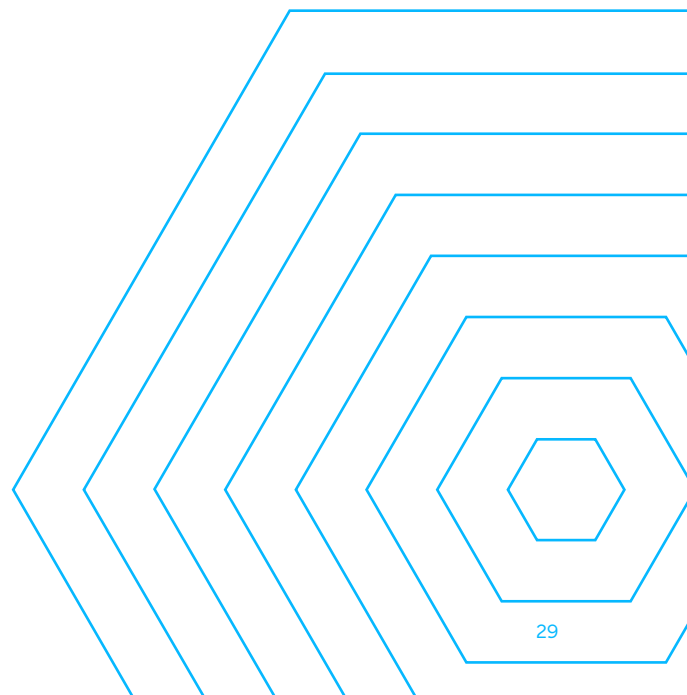
Limitationen

- **Kausalität:** Wir haben nicht alle Variablen untersucht, die bei der Startup-Finanzierung eine Rolle spielen. Die Einhaltung gesetzlicher Vorschriften ist nur ein Faktor und eine Risikoklassifizierung rechtfertigt für sich genommen nicht zwangsläufig eine höhere Finanzierung.



5. Schlussfolgerung

Diese Studie liefert die erste umfassende und empirisch fundierte Risikoklassifizierung des deutschen KI-Startup-Ökosystems nach der Verabschiedung der EU KI-VO. Unsere Ergebnisse zeigen, dass die regulatorischen Auswirkungen auf deutsche KI-Startups größer sind als erwartet: Rund 25 % der Startups sind potenziell mit hohem Risiko behaftet, weitere 26 % unterliegen möglicherweise Transparenzpflichten. Um zu verhindern, dass die KI-Verordnung Innovation hemmt, fordern wir die Interessengruppen auf EU-Ebene, in Deutschland und in der Industrie unter anderem dazu auf, die Leitlinien zu verbessern, die Regulierungskapazitäten auszubauen und die Einhaltung der Vorschriften zu vereinfachen. Wir hoffen, dass diese Erkenntnisse weitere Forschungen zu Compliance-Kosten und regulatorischen Auswirkungen für Startups anregen und Akteure im gesamten KI-Ökosystem dazu inspirieren, entschlossen Maßnahmen zur Beschleunigung vertrauenswürdiger KI-Innovationen in ganz Europa zu ergreifen. Eine fortlaufende datengestützte Analyse der KI-Verordnung, einschließlich ihrer Auswirkungen auf KI-Innovationen, wird auch weiterhin eine wesentliche Säule für evidenzbasierte Entscheidungsfindung sein, um die KI-Ziele Europas zu erreichen.



Anhang

Anhang 1: Datensatz

Diese Studie bezog alle KI-Startups ein, die Teil der „German AI Startup Landscape 2024“ sind (n=687⁸). 69 Startups aus der Landschaft wurden aufgrund inaktiver Websites, Liquidation, Übernahme oder Verlegung ihres Hauptsitzes aus Deutschland als nicht mehr teilnahmeberechtigt identifiziert und daher aus dieser Studie entfernt. Es gibt zwei Möglichkeiten, Teil der deutschen KI-Startup-Landschaft zu werden. Startups können sich entweder über unser Bewerbungsformular für die Aufnahme in die Landschaft bewerben oder von einem unserer Partner nominiert werden, bei denen es sich hauptsächlich um Wagniskapitalgeber und Technologieunternehmen handelt.

Um in die KI-Startup-Landschaft aufgenommen zu werden, muss jedes Startup die folgenden Kriterien erfüllen, die von unseren KI-Experten bewertet werden:

- Es muss ein eingetragenes Unternehmen sein.
- Es muss vor weniger als 10 Jahren gegründet worden sein oder sein Kerngeschäftsmodell vor weniger als 10 Jahren (einschließlich 2014) auf KI umgestellt haben.
- Der Hauptsitz des Startups muss sich in Deutschland befinden.
- Es sollte über mindestens zwei Vollzeitäquivalente (FTE) verfügen.
- Es sollte mindestens eine Vollzeitkraft mit KI-Kompetenz vorhanden sein.
- KI muss im Mittelpunkt stehen oder in erheblichem Umfang zum Einsatz kommen.
- Das Startup bzw. sein Geschäftsmodell hat auf den ersten Blick eine hohe Glaubwürdigkeit (z. B. professionelle Website, überzeugendes Geschäftsmodell usw.).

Für detaillierte Informationen zu den zugrunde liegenden Daten (z. B. in Bezug auf Branchen, zugrunde liegende Technologien und vieles mehr) verweisen wir auf die detaillierte Analyse der KI-Startup-Landschaft⁹. Darüber hinaus wurden detaillierte Finanzierungsdaten zu den KI-Startups von Dealroom bezogen und mit weiteren internen Daten angereichert.

8 Eine Handvoll Startups aus der „KI Startup Landscape 2024“ wurden für die Zwecke dieser Studie aus verschiedenen Gründen, wie z. B. Übernahme, Liquidation und anderen, ausgeschlossen.

9 KI-Startup-Landschaft 2024: <https://www.appliedai-institute.de/en/hub/2024-ai-german-startup-landscape>

Anhang 2: KI-Tool zur Risikoklassifizierung

Während der Datenerhebung wurden alle von den Startups angebotenen Anwendungsfälle gemäß den Risikokategorien der EU KI-VO klassifiziert: unannehmbares Risiko, Hohes Risiko, Transparenzverpflichtungen und minimales oder kein Risiko (weitere Informationen finden sich in Kapitel 2.1). Dementsprechend gingen wir in drei Schritten vor. Zunächst erfolgte die automatisierte Extraktion von Anwendungsfällen aus den Startups der deutschen KI-Landschaft. Anschließend erfolgte eine ensemblebasierte Klassifizierung unter Verwendung führender LLMs. Schließlich erfolgte eine manuelle Überprüfung der Risikoklassifizierungen durch Menschen.

Konkret haben wir ein automatisiertes und promptbasiertes Tool zur KI-Risikoklassifizierung entwickelt, das im ersten Schritt mithilfe der Web Search API von OpenAI Anwendungsfälle aller 628 für die Bewertung in Frage kommenden KI-Startups erfasste. In einem zweiten Schritt wurde jeder extrahierte Anwendungsfall mittels einer ensemblebasierten Klassifizierung mit LLMs einzeln analysiert und klassifiziert, wobei unser Tool für jede Entscheidung eine detaillierte Begründung lieferte. Der Klassifizierungsprozess, der im [GitHub repository](#) zu finden ist, wurde von unseren KI-Expertinnen und Experten sowie Juristinnen und Juristen entworfen. Sie fasst die KI-Verordnung in maschinenlesbarer Form zusammen und gibt klare Anweisungen für die automatisierte Risikobewertung. Ihre Erstellung war Teil eines iterativen Prozesses, der parallel zur technischen Entwicklung ablief (weitere Informationen siehe unten). In einem dritten Schritt führte unser Expertenteam eine manuelle Überprüfung jeder einzelnen Risikoklassifizierung durch (vgl. Insight-Box „Die Rolle der manuellen Überprüfung bei der Risikoklassifizierung“).

Natürlich bieten Startups oft mehr als eine Lösung an, d. h. sie haben mehrere Anwendungsfälle für ihre Lösung. Für die Zwecke dieser Studie geben wir die höchste Risikoklassifizierung jedes Startups an. Wenn ein Startup beispielsweise sowohl ein Tool mit minimalem oder keinem Risiko zur Dokumentenzusammenfassung als auch ein System zur prädiktiven Gesundheitsanalyse anbot, wählten wir das Startup anhand der Kategorie mit dem höheren Risiko aus. Dieser Ansatz stellt sicher, dass wir die kritischsten regulatorischen Compliance-Aspekte berücksichtigen, die mit dem Angebotsportfolio jedes Startups verbunden sind. Um die Qualität der automatisierten Kategorisierung zu gewährleisten, haben unsere Expertinnen und Experten sämtliche Einträge gegengeprüft und die Ergebnisse manuell validiert.

Methode

Es wurde eine Anwendung in Python entwickelt, um KI-Anwendungsfälle der Startups in der „AI Startup Landscape“ zu extrahieren und jeden identifizierten Anwendungsfall gemäß den in der definierten Regeln zu klassifizieren. Der Code ist auf [GitHub](#) öffentlich zugänglich. Die architektonische Übersicht des Tools ist in der folgenden Abbildung (**Abbildung 11**) dargestellt:

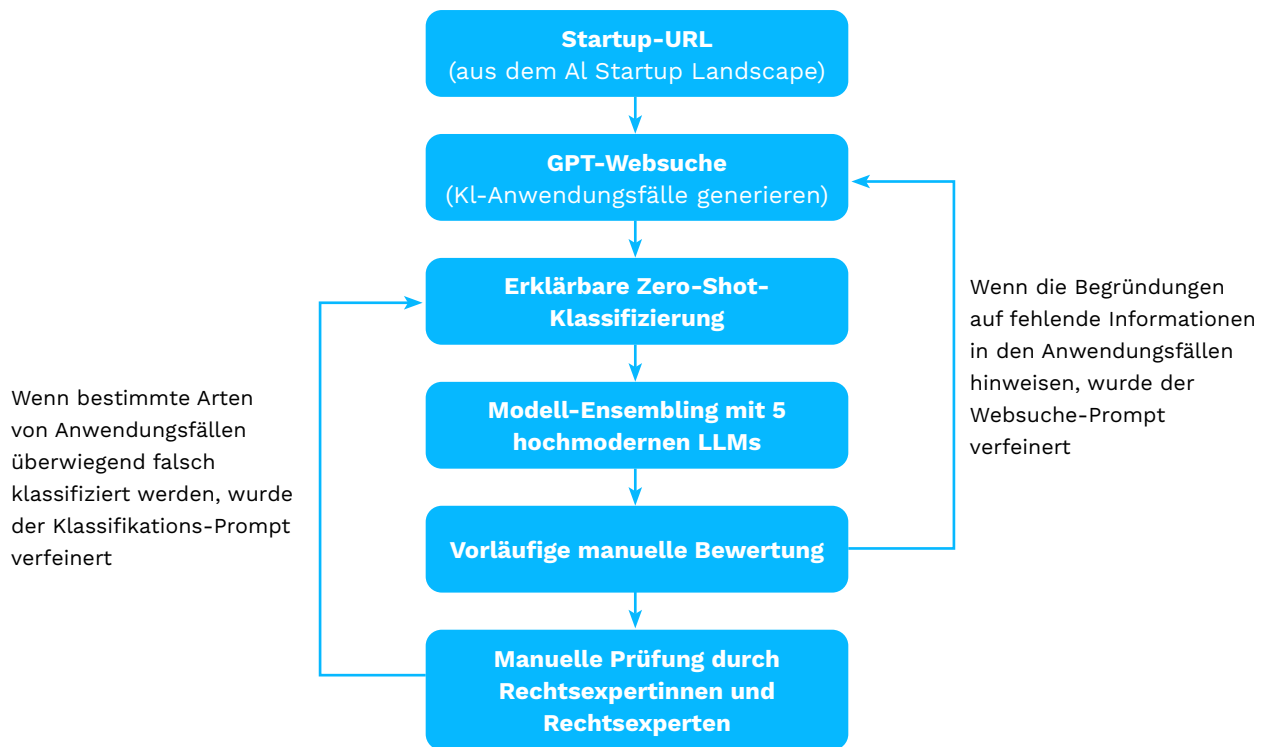


Abbildung 11: Anwendungsarchitektur

Der erste Schritt bestand darin, alle KI-Anwendungsfälle der Startups zu extrahieren. Zum Zeitpunkt der Durchführung dieser Studie veröffentlichte OpenAI eine „Web Search“-API für eine begrenzte Anzahl von Modellen, insbesondere „gpt-4o-search-preview“ und „gpt-4o-mini-search-preview“. Während das erstgenannte Modell größer ist, steuert der sogenannte Parameter „search context size“ in der API, wie viel Inhalt aus dem Web abgerufen wird. Die API bietet für diesen Parameter drei Werte: „high“, „medium“ und „low“, in absteigender Reihenfolge der Suchkontextgröße. Für unsere Zwecke entschieden wir uns für das Modell „gpt-4o-search-preview“ und wählten die höchste Suchkontextgröße, damit die maximale Anzahl an Webseiten gecrawlt werden konnte, um die bestmöglichen Anwendungsfälle zu erhalten. Zwar ist eine höhere Suchkontextgröße mit höheren Kosten verbunden, doch war der Kompromiss zugunsten der Qualität gerechtfertigt.

Alle KI-Anwendungsfälle für jedes Startup wurden in einer CSV-Datei gespeichert. Die Web Search API wurde angewiesen, die relevantesten Informationen für unsere Zero-Shot-Klassifizierungsaufgabe zu sammeln. Dies umfasste die folgenden Details in den generierten Anwendungsfällen:

1. Verwendungszweck
2. Einsatzbereich
3. Autonomiestufe
4. Auswirkungen auf Einzelpersonen
5. Verwendete Datentypen
6. Benutzertyp
7. Anpassungsfähigkeit bei der Bereitstellung
8. Sicherheitskritischer Charakter des Systems

Modell-Ensembling mit Mehrheitsentscheidung

Nachdem alle KI-Anwendungsfälle für jedes Startup erfasst worden waren, bestand der nächste Schritt darin, jeden Anwendungsfall in die verschiedenen Kategorien der EU KI-VO einzuordnen. Der von uns verwendete Ansatz wird als „erklärbare Zero-Shot-Klassifizierung“ bezeichnet. Der Begriff „erklärbar“ bezieht sich darauf, dass die Modelle dazu veranlasst werden, zusammen mit jedem Klassifizierungslabel eine Begründung zu generieren¹⁰.

Die Kategorien für diesen Zweck wurden wie folgt definiert:

1. Verbotenes KI-System
2. Hochrisiko-KI-System gemäß Anhang I
3. Hochrisiko-KI-System gemäß Anhang III
4. Hochrisiko-KI-System mit Transparenzpflichten
5. System mit Transparenzpflichten
6. KI-System mit minimalem oder keinem Risiko
7. Unsicher

Die [Klassifizierungsanweisung](#) enthielt die genauen Kriterien für die Zuordnung einer Risikokategorie aus der obigen Liste zu jedem generierten KI-Anwendungsfall. Die Anzahl der Eingabetoken lag deutlich innerhalb der maximalen Kontextlängen aller verwendeten Modelle.

Ensembling mit LLMs

Bei herkömmlichen Aufgaben des Machine Learning / Maschinellen Lernens ist das Ensembling eine Methode, bei der mehrere Modelle gleichzeitig verwendet werden, um eine Klassifizierung vorherzusagen. Zwar gibt es verschiedene Arten von Ensembling-Methoden wie Bagging, Boosting usw., doch das ultimative Ziel des Ensemblings besteht darin, die Vorhersagen mehrerer Modelle zu einem einzigen Ergebnis zu kombinieren. Im Allgemeinen lässt sich dies auf verschiedene Weise erreichen; im Folgenden sind einige Beispiele aufgeführt:

1. Ermittlung durch Mehrheitsentscheidung unter allen vorhergesagten Labels
2. Berechnung des Durchschnitts aller Modellausgabewerte. Dies funktioniert jedoch nur bei numerischen Daten.
3. Ermittlung durch eine gewichtete Abstimmung oder einen gewichteten Durchschnittswert. Bestimmte Modelle, von denen bekannt ist, dass sie eine bessere Leistung erbringen (d. h. eine höhere Validierungsgenauigkeit aufweisen), erhalten ein höheres Gewicht.
4. Auswahl der Klassifizierung, die ein beliebiges Modell vorhersagt (Max-Voting), oder nur der Klassifizierung, auf die sich alle Modelle einigen (Min-Voting)

Für unsere Aufgabe haben wir den ersten Ansatz des Ensemble-Verfahrens mit Mehrheitsentscheidung verwendet. Für diesen Ansatz haben wir die folgenden 5 modernsten LLMs ausgewählt:

1. ChatGPT 4o
2. Claude 3.7 Sonnet
3. DeepSeek Reasoner
4. Gemini 2.0 Flash Thinker
5. Mistral Large

¹⁰ Zur Terminologie: Eine Klasse repräsentiert die Menge aller möglichen Werte für eine bestimmte Klassifizierung, während ein Label den spezifischen Wert bezeichnet, den das Modell vorhersagt.

Jeder KI-Anwendungsfall wurde zusammen mit der Klassifizierungsaufforderung über die entsprechende API gleichzeitig in jedes KI-Modell eingespeist. Sobald alle Antworten der Modelle eingegangen sind, startet der Abstimmungsmechanismus. In diesem Prozess suchen wir nach dem häufigsten Klassifizierungslabel (oder Risikostufe) unter allen von den Modellen generierten vorhergesagten Labels. Das folgende Diagramm (**Abbildung 12**) veranschaulicht den Abstimmungsmechanismus:

Leistung

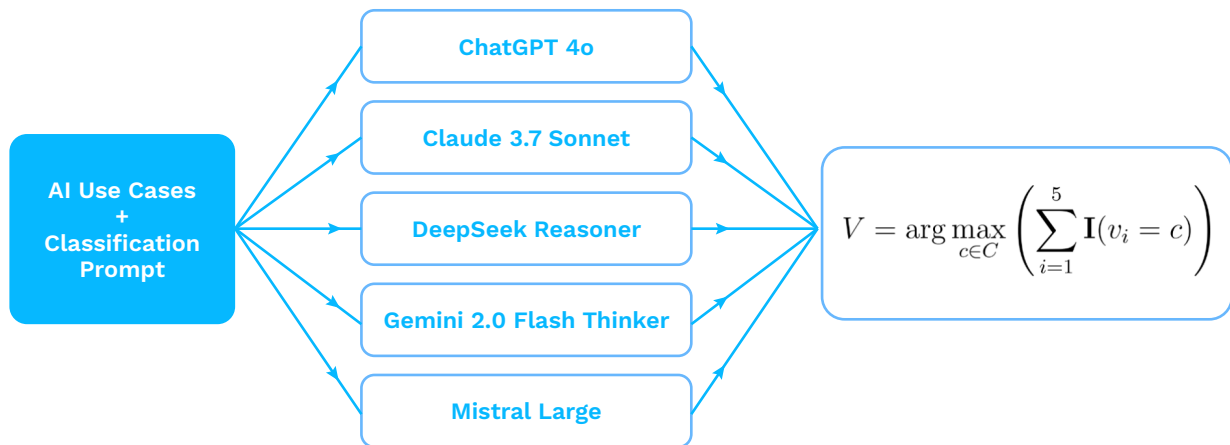


Abbildung 12: Modell-Ensambling

Die Eingabe beginnt mit einem einzelnen Anwendungsfall und der Klassifizierungsaufforderung, die gleichzeitig in jedes Modell eingespeist wird. Jedes Modell generiert eine Antwort, die die folgenden Felder enthält:

- 1. Risikoklassifizierung:** Das Klassifizierungslabel für den Anwendungsfall
- 2. Begründung:** Eine umfassende Begründung für die Kennzeichnung, die den Anwendungsfall gemäß allen Anweisungen in der Klassifizierungsaufforderung bewertet. Diese Begründung ist in 7 separate Schritte unterteilt, die den Anwendungsfall anhand jeder möglichen Risikoklasse bewerten (Verboten, Hochrisiko-KI-System gemäß Anhang I ... KI-System mit minimalem oder keinem Risiko)
- 3. Zusätzliche Informationen erforderlich:** Wenn das Modell nicht über ausreichende Informationen verfügt, um eine angemessene Klassifizierung vorzunehmen, wird der Wert dieses Feldes auf „Ja“ gesetzt
- 4. Welche zusätzlichen Informationen:** Falls zusätzliche Informationen erforderlich waren, werden diese in diesem Feld angegeben

Das Programm wartet darauf, dass jedes Modell seinen Ergebnissatz generiert. Das Feld „Risikoklassifizierung“ jedes Modells wird dann wie oben gezeigt an eine argmax-Funktion übergeben. Die Funktion zählt lediglich die Anzahl der von jedem Modell generierten Klassifizierungsbezeichnungen. Diese Zählung wird als „Stimme“ bezeichnet. Je höher die Anzahl der Stimmen ist, die eine Bezeichnung erhält, desto höher ist die Wahrscheinlichkeit, dass diese Bezeichnung dem Anwendungsfall zugewiesen wird.

Betrachten wir beispielsweise den folgenden Fall:

- C als die Menge aller möglichen Klassen
- Menge aller von den LLMs abgegebenen Stimmen (v_i):
 - Hochrisiko-KI-System gemäß Anhang I
 - Hochrisiko-KI-System gemäß Anhang I
 - KI-System mit geringem Risiko
 - Hochrisiko-KI-System gemäß Anhang I
 - KI-System mit geringem Risiko
- Für jedes $c \in C$, berechnen wir Folgendes, um die Mehrheitsentscheidung V zu ermitteln:

for $c = \text{"High-risk...Annex I"}: \sum_{i=1}^5 \mathbb{I}(v_i = \text{"High-risk...Annex I"}) = 3$

for $c = \text{"Low-risk AI system"}: \sum_{i=1}^5 \mathbb{I}(v_i = \text{"Low-risk AI system"}) = 2$

- Für alle anderen Klassen in C , die von keinem der Modelle zugeordnet wurden, wird die Stimmenzahl auf 0 gesetzt. Beispielsweise erscheint „System mit Transparenzpflichten“ in keinem der Stimmen v_i . Daher beträgt die Stimmenzahl für diese Klasse 0.
- Das endgültige Mehrheitsergebnis für dieses Beispiel lautet:

$V = \text{"High-risk...Annex I"}$

Diese Klasse hat die höhere Stimmenzahl (d. h. 3) und eine klare Mehrheit. Daher wird sie als Klassifizierungslabel für den Anwendungsfall zugewiesen.

Sobald die Bezeichnung zugewiesen ist, müssen wir eine der drei vom Modell generierten Antworten auswählen, die im obigen Beispiel die Stimmenmehrheit erhalten haben. Zu diesem Zeitpunkt sind alle drei Modellantworten korrekt. Daher wählen wir einfach die Modellantwort, die als erste eingegangen ist, als endgültige richtige Antwort aus. Der ausgewählte Modellname sowie alle anderen Details aus der beschriebenen Antwort werden in der CSV-Ausgabedatei gespeichert.

Bei Stimmengleichheit wird die Klassifizierung mit dem geringeren Risiko, entsprechend der Anordnung in [dieser](#) Liste, ausgewählt. Beispielsweise bei den folgenden Stimmen v_i :

$$\sum_{i=1}^5 \mathbb{I}(v_i = \text{"High-risk...Annex I"}) = 2$$

$$\sum_{i=1}^5 \mathbb{I}(v_i = \text{"High-risk...Annex III"}) = 2$$

$$\sum_{i=1}^5 \mathbb{I}(v_i = \text{"Low-risk AI system"}) = 1$$

Es liegt ein Gleichstand zwischen „Hochrisiko-KI-System gemäß Anhang I“ und „Hochrisiko-KI-System gemäß Anhang III“ vor. Der Abstimmungsmechanismus wählt konservativ das geringere Risiko aus den Ergebnissen aus, d. h. das Hochrisiko-KI-System gemäß Anhang III.

Wir haben die Leistung des Tools anhand einer Konfusionsmatrix bewertet, die in **Abbildung 13** dargestellt ist.

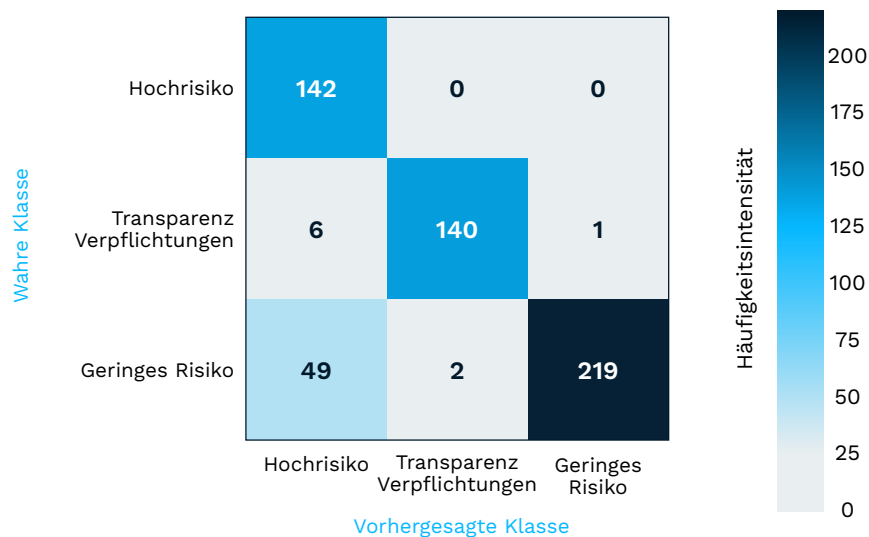


Abbildung 13: Konfusionsmatrix

Die Gesamtgenauigkeit des Tools lässt sich wie folgt berechnen:

Genauigkeit = Anzahl der korrekten Vorhersagen / Gesamtzahl der Vorhersagen

$$= (142 + 140 + 219) / (142 + 0 + 0 + 6 + 140 + 1 + 49 + 2 + 219)$$

$$= 501 / 559$$

$$= \mathbf{89.6\%}$$

Somit beträgt die Gesamtgenauigkeit des Tools 89,6 %.

Bei der Einstufung in sechs regulatorische Risikoklassen (für die Analyse auf drei zusammengefasst) zeigt das Tool eine hohe Gesamtgenauigkeit, neigt jedoch zur Vorsicht, indem es Startups mit minimalem oder keinem Risiko gelegentlich als Hochrisiko einstuft.

Zusätzliche manuelle Überprüfung

Die Risikoklassifizierungen in diesem Datensatz wurden durch einen hybriden Arbeitsablauf erstellt: Die Ensemble-LLMs lieferten eine erste Klassifizierung für jeden der über 600 Anwendungsfälle, woraufhin zwei erfahrene menschliche Prüfer die Ergebnisse unabhängig voneinander überprüften. Wenn beide Prüfer die vom System zugewiesene Einstufung akzeptierten, wurde diese als endgültig festgehalten. Wenn einer oder beide nicht zustimmten, führten sie eine Ad-hoc-Diskussion, wählten die „vernünftigste“ Klasse aus und – falls sie der Meinung waren, dass die Modelle falsch lagen – überschrieben sie den Eintrag. Wir haben uns gegen strengere Methoden entschieden, da das Ziel der Studie eine explorative Kartierung und nicht eine behördliche Zertifizierung war.

Referenzen

European Commission. (2021). *Commission staff working document: Impact assessment accompanying the proposal for a regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts (SWD(2021) 85 final)*. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=SWD:2021:85:FIN>

Hauer, M. P., Krafft, T. D., Sesing-Wagenpfeil, A., & Zweig, K. (2023). *Quantitative study about the estimated impact of the AI Act*. arXiv. <https://arxiv.org/abs/2304.06503>

appliedAI. (2023, March). *AI Act: Risk classification of AI systems from a practical perspective*. appliedAI Institute for Europe. <https://www.appliedai.de/en/insights/ai-act-risk-classification-of-ai-systems-from-a-practical-perspective/>

OECD. (2024). *OECD artificial intelligence review of Germany*. OECD Publishing. <https://doi.org/10.1787/609808d6-en>

Hutchinson, P., Goll, F., & Mügge, C. (2024). *Generative AI in the European Startup Landscape 2024*. appliedAI Institute for Europe. <https://www.appliedai-institute.de/en/hub/2024-generative-ai-study>

Renda, A., Fanni, R., Laurer, M., Sipiczki, A., Yeung, T., Maridis, G., Fernandes, M., Endrodi, G., Milio, S., Devenyi, V., Georgiev, S., de Pierrefeu, G., & Arroyo, J. (2021). *Study to support an impact assessment of regulatory requirements for artificial intelligence in Europe: Final report*. European Commission, Directorate-General for Communications Networks, Content and Technology. <https://doi.org/10.2759/523404>

Fetic, L., Niemeyer, D., & Klein, T. (2025). *Wie können KI-Reallabore ihr Potenzial in der KI-Verordnung entfalten? Bedingungen für ein wirksames Instrument eines souveränen und lernfähigen KI-Ökosystems in Europa*. appliedAI Institute for Europe. <https://www.appliedai-institute.de/media/downloads/Policy-Brief-AI-Regulatory-Sandbox.pdf?v=1764343275>

Laurer, M., Renda, A., & Yeung, T. (2021, September 24). *Clarifying the costs for the EU's AI Act*. CEPS. <https://www.ceps.eu/clarifying-the-costs-for-the-eus-ai-act/>

Mueller, B. (2021, July). *How much will the Artificial Intelligence Act cost Europe?* Center for Data Innovation. <https://datainnovation.org/2021/07/how-much-will-the-artificial-intelligence-act-cost-europe/>

Über die Autoren

Alle Autoren haben gleichermaßen zu dieser Arbeit beigetragen.



Dr. Philip Hutchinson  

Dr. Philip Hutchinson ist Senior AI Strategist beim gemeinnützigen appliedAI Institute for Europe. Darüber hinaus bekleidet er verschiedene Expertenpositionen, u. a. beim EuroHPC Joint Undertaking und bei der Europäischen Kommission. Zuvor war Dr. Hutchinson als Unternehmensberater bei Ernst & Young tätig, wo er Transformationsprojekte in Großunternehmen begleitete. Er promovierte an der Universität Kiel im Fach Innovationsmanagement mit dem Forschungsschwerpunkt auf der Anwendbarkeit von Künstlicher Intelligenz im Innovationsprozess.



Dr. Till Klein 

Dr. Till Klein ist Leiter des Bereichs KI-Regulierung beim gemeinnützigen appliedAI Institute for Europe und Projektleiter des Bayerischen KI-Innovationsbeschleunigers. Till sammelte mehrjährige Branchenerfahrung im Bereich Regulatory Affairs für Medizinprodukte, als leitender Auditor für Qualitätsmanagementsysteme sowie als Leiter des Qualitätsmanagements in einem Drohnenunternehmen, wodurch er über praktische Fähigkeiten und Perspektiven verfügt, die für die Anwendung vertrauenswürdiger KI von entscheidender Bedeutung sind. Er hat einen Dokortitel in Betriebswirtschaft vom Centre for Transformative Innovation an der Swinburne University of Technology in Melbourne, Australien, sowie einen Bachelor-Abschluss in Wirtschaftsingenieurwesen.



Shahrukh Azhar Ahsan 

Shahrukh Azhar Ahsan ist Applied AI Engineer beim gemeinnützigen appliedAI Institute for Europe und verfügt über Erfahrung in der Implementierung von KI-Systemen für private und staatliche Kunden. Er ist spezialisiert auf die Entwicklung von KI-Agenten, RAG-Systemen und hochoptimierten Inferenz-Pipelines für Open-Source-LLMs und VLMS. Bevor er sich auf KI konzentrierte, arbeitete er bei Startups in den USA, Großbritannien und Kanada an der Entwicklung von DTC-E-Commerce-Plattformen und sprachgesteuerten Webanwendungen. Außerdem war er Mitbegründer eines EdTech-AR-Startups für naturwissenschaftliche Labore an weiterführenden Schulen. Er hat einen M.Sc. in Künstlicher Intelligenz und Datenwissenschaft von der Fachhochschule Deggendorf in Deutschland und einen B.Sc. in Informatik von der National University of Sciences & Technology (NUST) in Islamabad, Pakistan.



Akhil Deo  

Akhil ist Senior AI Regulatory Affairs Expert beim gemeinnützigen appliedAI Institute for Europe, wo er Unternehmen bei der KI-Governance und der Produktrichtliniengestaltung unterstützt. Zuvor war er als Kommunikationsspezialist am Future of Life Institute und als wissenschaftlicher Mitarbeiter bei der Observer Research Foundation tätig. Akhil hat einen Master-Abschluss in Internationalen Beziehungen von der Hertie School.

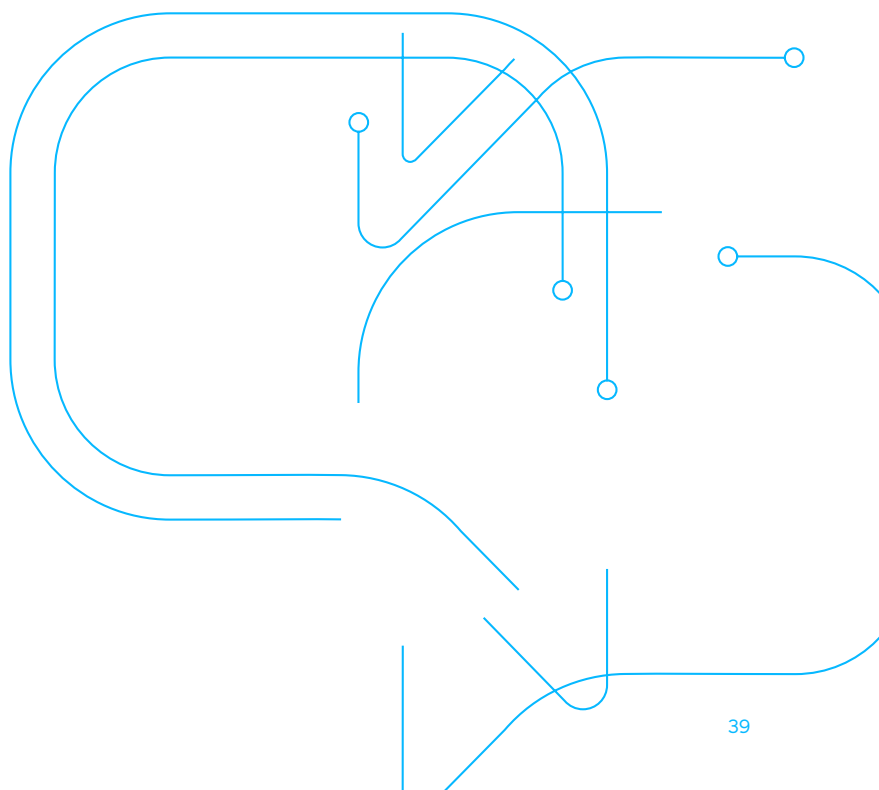
Wir als Autoren möchten Dr. Wolfgang Fischer unseren aufrichtigen Dank für seine wertvolle inhaltliche Beratung und seine aufschlussreichen Beiträge zu dieser Publikation aussprechen. Sein Fachwissen und seine durchdachten Anregungen haben die Qualität und Tiefe unserer Arbeit erheblich bereichert.

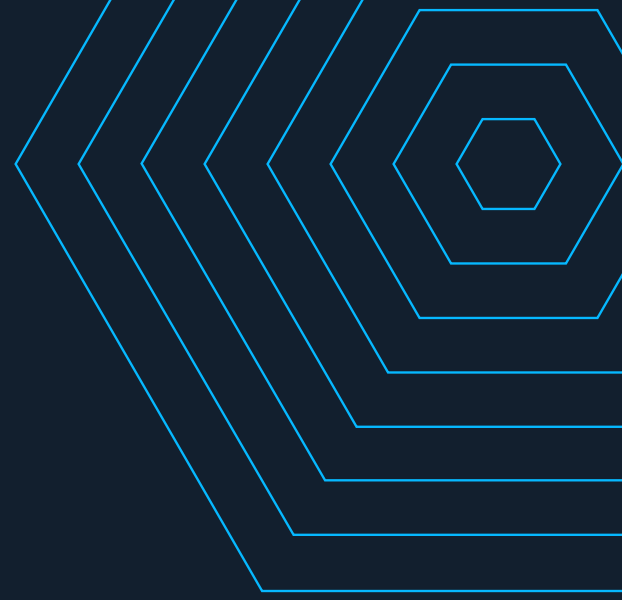
Herausgeber

appliedAI Institute for Europe gGmbH
Freddie-Mercury-Straße 5
D-80797 München

Lizenz

Diese Studie wird unter der Creative-Commons-Lizenz CC BY-SA 4.0 veröffentlicht.





appliedAI Institute for Europe – Über uns

Das appliedAI Institute for Europe hat sich zum Ziel gesetzt, das europäische KI-Ökosystem zu stärken, Forschung im Bereich KI voranzutreiben, Wissen rund um KI zu entwickeln, vertrauenswürdige KI-Tools bereitzustellen und Bildungs- sowie Interaktionsformate rund um hochwertige KI-Inhalte zu schaffen.

Als gemeinnützige Tochtergesellschaft der appliedAI Initiative wurde das Institut 2022 in München gegründet. Die appliedAI Initiative selbst ist ein Joint Venture aus UnternehmerTUM und IPAI. Die Leitung des Instituts obliegt Dr. Andreas Liebl.

Das appliedAI Institute for Europe stellt die Menschen in Europa in den Mittelpunkt. Es verfolgt die Vision, eine gemeinsame KI-Community zu formen und hochwertige Inhalte im Zeitalter der KI für die gesamte Gesellschaft bereitzustellen. Durch die Förderung von vertrauenswürdiger KI beschleunigt das Institut die Anwendung dieser Technologie und stärkt Vertrauen in KI-Lösungen.

Mit einem Fokus auf Wissensentwicklung, Forschung und der Bereitstellung vertrauenswürdiger KI-Tools bietet das appliedAI Institute for Europe eine wertvolle Ressource für Unternehmen, Organisationen und Einzelpersonen, die ihre Kenntnisse und Fähigkeiten im Bereich KI erweitern möchten. Durch Bildungs- und Interaktionsformate ermöglicht das Institut einen intensiven Austausch von Expertise und fördert die Zusammenarbeit zwischen Akteuren aus verschiedenen Bereichen.

Das appliedAI Institute for Europe lädt Unternehmen, Organisationen, Startups und KI-Enthusiast:innen ein, von den vielfältigen Angeboten und Ressourcen des Instituts zu profitieren. Die appliedAI Institute for Europe gGmbH wird unterstützt durch die KI-Stiftung Heilbronn gGmbH.

Weitere Informationen finden Sie unter www.appliedai-institute.de.



Gefördert durch
Bayerisches Staatsministerium
für Digitales



Diese Studie wurde durch das Projekt „Bayerischer KI-Innovationsbeschleuniger“ unterstützt.

Der Bayerische KI-Innovationsbeschleuniger ist ein zweijähriges Projekt, gefördert vom Bayerischen Staatsministerium für Digitales, um KMU, Startups und den öffentlichen Sektor in Bayern bei der Einhaltung der EU-KI-Verordnung zu unterstützen. Unter der Leitung des appliedAI Institute for Europe und in Zusammenarbeit mit der Ludwig-Maximilians-Universität, der Technischen Universität München und der Technischen Universität Nürnberg werden Trainings, Ressourcen und Veranstaltungen angeboten. Die Projektziele umfassen die Senkung der Konformitätskosten, die Verkürzung der Zeit bis zur Konformität und die Stärkung von KI-Innovation. Zur Umsetzung dieser Ziele ist das Projekt in fünf Arbeitspakete gegliedert: Projektmanagement, Forschung, Bildung, Tools & Infrastruktur sowie Community.

<https://innovationsbeschleuniger.bayern>



Gefördert durch
Bayerisches Staatsministerium
für Digitales



Bayerischer
KI-Innovationsbeschleuniger

appliedAI Institute for Europe

Freddie-Mercury-Straße 5

80797 München

Germany

www.appliedai-institute.de