



Verhandeln bei M&A-Transaktionen mit AI-Unterstützung

19. März 2026

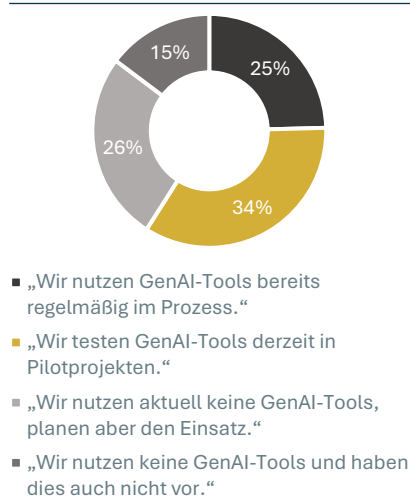
Verhandeln bei M&A-Transaktionen mit AI-Unterstützung	2
1. Executive Summary	2
2. Modelllandschaft und Leistungsprofil aktueller LLM – Taxonomie relevanter AI-Ansätze	3
2.1 AI“ ist kein einheitliches System, sondern ein Sammelbegriff	3
2.2 Die wichtigsten Systemklassen im M&A-Kontext	3
2.3 Warum diese Unterscheidung für M&A so wichtig ist	5
2.4 Was die EU-AI-Verordnung daran ändert	5
2.5 Provider- und Deployer-Logik	6
2.6 Warum das strategisch relevant ist	6
3. Was LLM heute gut können	7
3.1 Wie funktionieren LLM eigentlich?	7
3.2 Warum wirken LLM „hilfsbereit“ und „kooperativ“?	7
3.3	8
3.4 Was bedeutet das konkret für M&A-Verhandlungen?	8
4. Wo die strukturellen Grenzen liegen	10
4.1 Kooperations- bzw. Wohlwollens-Bias von Chat-LLM	10
4.2 Rechen- und Bewertungszuverlässigkeit	11
4.3 Zentrale Risikoklassen nach NIST	12
5. Verhandeln bei M&A: Zwei-Ebenen-Logik und AI-Fit	13
6. Evidenz: Forschungsstand und ALLERTS-MONITOR-Ergebnisse	14
6.1 Forschungsstand zur AI-Negotiation	14
6.2 Umfragebefunde: Adoption und Wahrnehmung	15
6.3 Qualitative Befunde: „größte Herausforderung“	16
7. Use-Case-Matrix und Zielarchitektur	17
8. Governance & Compliance in DACH/EU sowie Best-Practice-Playbook – Regulatorischer und normativer Rahmen	18
9. Best-Practice-Playbook für M&A-Verhandlungsteams – Der VIIV-Use-Case:	19
Die 20 Key-Findings aus den vorliegenden wissenschaftlichen Arbeiten im Quick-View (AI-generiert)	21
Glossar	24
Quellen- und Literaturverzeichnis	29

Verhandeln bei M&A-Transaktionen mit AI-Unterstützung

1. Executive Summary

Künstliche Intelligenz (AI) – insbesondere generative AI (GenAI) und Large Language Models (LLM)¹ – verändert den M&A-Markt nicht nur punktuell, sondern strukturell: sie beschleunigt Sprach- und Dokumentenarbeit, senkt Such- und Strukturierungskosten und vergrößert die Bandbreite an in kurzer Zeit verarbeitbaren Informationen. Gleichzeitig ist sie keine Verhandlungsmaschine im Sinne eines „Deal-Autopiloten“: LLM sind probabilistische Sprachmodelle, die in kompetitiven Verhandlungssituationen an Grenzen stoßen, insbesondere dort, wo Durchsetzungsfähigkeit, Machtasymmetrien und psychologische/strategische Dynamik dominieren. Diese Unterscheidung ist zentral – und entspricht auch der kommentierten Leitlinie, dass AI Vorteile vor allem dann realisiert, wenn Verhandlungen kooperativ/faktenorientiert geführt werden; in kompetitiven Settings entscheiden häufig nicht „bessere Argumente“, sondern Durchsetzung und Taktik.

Abb.1: Einsatz von GenAI-Tools im M&A-Verhandlungsprozess



Quelle: ALLERTS MONITOR 02/2025

Die in dieser Studie ausgewertete Umfrage unter M&A-Transaktionsprofessionals² zeigt ein typisches Transformationsmuster: Adoption ist bereits spürbar, Skepsis in Bezug auf Qualität, Effizienzbeitrag und Governance-Reife jedoch ebenso. Konkret geben 24,6% an, GenAI bereits regelmäßig im M&A-Verhandlungsprozess zu nutzen; 34,4% testen in Pilotprojekten; 26,2% planen den Einsatz; 14,8% nutzen nicht und planen dies auch nicht. Gleichzeitig erwarten die Befragten mehrheitlich, dass GenAI-Tools in drei Jahren Standard sein werden (Mittelwert 3,68 auf einer 5-stufigen Zustimmungsskala). Dieses Spannungsfeld – „kommt sicher“ bei gleichzeitig „noch nicht reif/noch nicht sauber geregelt“ – ist die strategische Kernlage.

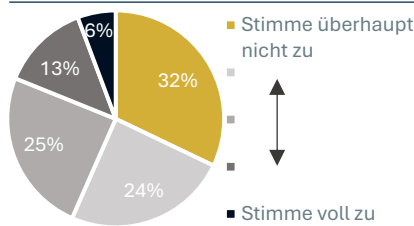
Aus diesen Daten, dem kommentierten Arbeitsdokument und der Forschungsliteratur folgt:

1. für High-Stakes-M&A gilt: „Der Einsatz von AI als Co-Pilot, verstanden als unterstützendes Entscheidungs- und Analyseinstrument unter menschlicher Letztverantwortung, bildet aktuell den belastbarsten und regulatorisch wie haftungsrechtlich verantwortbaren Standard für professionelle M&A-Transaktionen.“
2. Organisationen müssen strikt trennen zwischen Kommunikations-/Textarbeit (LLM-Stärken) und ökonomischem Rechnen/Entscheiden (deterministisch zu verifizieren).

¹ Glossar im Anhang

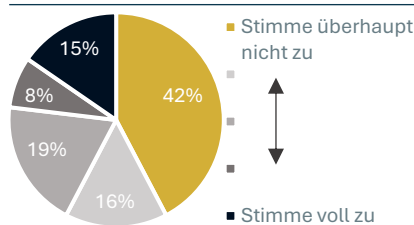
² In Kooperation mit Allert & Co. GmbH

Abb.2: Einschätzung der Effizienzsteigerung durch GenAI-Tools im Verhandlungsprozess



Quelle: ALLERTS MONITOR 02/2025

Abb.3: Einführung interner Richtlinien zur GenAI-Nutzung



Quelle: ALLERTS MONITOR 02/2025

- LLM können Zahlen „sprachlich plausibilisieren“, aber nicht zuverlässig wie ein prüfbarer Rechenkern arbeiten.
3. Governance entscheidet über Nutzen oder Schaden: NIST³ beschreibt für GenAI u. a. Risiken wie Confabulation/Halluzinationen, Prompt-Injection und Datenabfluss als systematische Risikoklassen; ohne Kontrollen steigt die Fehler- und Haftungswahrscheinlichkeit.⁴

Die VIIV-Studie liefert dafür ein umsetzbares Zielbild:

- eine Use-Case-Matrix für den gesamten M&A-Verhandlungsprozess
- ein Governance-/Compliance-Operating-Model für DACH/EU (EU-AI-Act- und NIST-RMF-konform) sowie
- ein Best-Practice-Playbook

Die zentrale Empfehlung lautet: LLM + Retrieval-Augmented Generation (RAG) + deterministische Tools + Human-in-the-Loop als aktueller Best-Practice-Standard, jedoch keinesfalls ein Verlassen auf „LLM-only“.⁵

2. Modelllandschaft und Leistungsprofil aktueller LLM – Taxonomie relevanter AI-Ansätze

2.1 AI“ ist kein einheitliches System, sondern ein Sammelbegriff

Wenn wir im Zusammenhang mit M&A-Verhandlungen von AI sprechen, klingt das oft so, als gäbe es eine einzige Art von künstlicher Intelligenz, die alles kann. Das ist nicht der Fall. Tatsächlich handelt es sich um verschiedene Arten von Systemen, die sehr unterschiedlich funktionieren und unterschiedliche Aufgaben übernehmen. Für M&A-Transaktionen ist es entscheidend, diese Unterschiede zu verstehen, weil jede Systemklasse andere Stärken, Schwächen und Risiken hat. Man kann sich das wie einen Werkzeugkasten vorstellen: Nicht jedes Werkzeug ist ein Hammer. Und nicht jede AI ist ein Sprachmodell.

2.2 Die wichtigsten Systemklassen im M&A-Kontext

Symbolisch-regelbasierte Systeme
(Policies, Workflows, Entscheidungsregeln)

³ NIST steht für National Institute of Standards and Technology. Es handelt sich um eine US-amerikanische Bundesbehörde (Teil des U.S. Department of Commerce), die technische Standards, Sicherheitsrahmenwerke und Bewertungsmodelle entwickelt. NIST ist weltweit ein Referenzgeber für IT-, Cyber- und zunehmend auch KI-Governance.

⁴ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

⁵ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

Das sind die klassischen digitalen Systeme, die nach klaren Wenn-Dann-Regeln arbeiten.

Beispiel:

- „Wenn Kaufpreis > X und Earn-out enthalten → zusätzliche Prüfung erforderlich.“
- „Wenn im SPA bestimmte Garantien fehlen → Warnhinweis.“

Solche Systeme:

- denken nicht selbst
- lernen nicht
- erzeugen keine neuen Inhalte
- sondern folgen festen Regeln

Bei M&A-Verhandlungen können Sie genutzt werden für:

- Freigabeprozesse
- Compliance-Prüfungen
- Genehmigungs-Workflows
- strukturierte Checklisten

Solche Tools sind mittlerweile sehr zuverlässig – aber nicht kreativ.

Klassisches Machine Learning (ML) (Scoring- und Prognosemodelle)

Diese Systeme lernen – sofern relevante, vergleichbare Daten in großer Menge vorliegen – aus Informationen und berechnen Wahrscheinlichkeiten.

Beispiele:

- Ausfallwahrscheinlichkeiten
- Risikoscores
- Bewertungsprognosen
- Erfolgswahrscheinlichkeiten von Transaktionen

Ein solches Modell kann z.B. berechnen:

- Wie wahrscheinlich ist ein Deal-Abbruch?
- Wie wahrscheinlich sind welche Finanzierungskonditionen?
- Wie hoch ist das statistische Risiko bestimmter Vertragsklauseln?

Wichtig: Diese Systeme erzeugen keinen Text, sondern liefern Zahlen, Wahrscheinlichkeiten oder Klassifizierungen.

Deep Learning (Komplexe neuronale Netze)

Deep Learning ist eine Weiterentwicklung des klassischen Machine Learning. Es arbeitet mit sehr großen Datenmengen und komplexen mathematischen Strukturen.

Diese Systeme können:

- Muster in riesigen Dokumentenmengen erkennen
- Klauseltypen automatisch klassifizieren
- Vertragsabweichungen erkennen

Im M&A-Kontext sind sie relevant für z.B.:

- automatisierte Dokumentenanalyse (Due Diligence)
- Mustererkennung in Vertragswerken
- Risikoidentifikation

Generative Foundation Models (z.B. LLMs)

Diese Systeme bekommen aktuell am meisten Aufmerksamkeit.

Ein Large Language Model (LLM) kann (bzw. könnte theoretisch):

- Texte formulieren
- Argumentationslinien entwickeln
- Vertragsentwürfe umschreiben
- Verhandlungsargumente strukturieren
- Zusammenfassungen erstellen

Wichtig zu wissen: Ein LLM versteht nicht wie ein Mensch, sondern berechnet statistisch wahrscheinliche Wortfolgen. Es ist ein Sprachgenerator, kein strategischer Entscheider.

2.3 Warum diese Unterscheidung für M&A so wichtig ist

In einer M&A-Verhandlung laufen zwei Prozesse gleichzeitig:

1. ein sprachlicher Prozess (Argumente, Formulierungen, Taktik)
2. ein ökonomisch-rechnerischer Prozess (Bewertung, Szenarien, Preislogik)

Ein LLM kann beim sprachlichen Prozess stark unterstützen. Für präzise finanzmathematische Berechnungen ist es dagegen nicht ausreichend zuverlässig.

Deshalb braucht es in der Praxis eine Kombination:

- LLM für Sprache
- deterministische Rechentools für Zahlen
- regelbasierte Systeme für Governance
- ML-Modelle für Risikobewertung

AI ist also kein einzelnes System, sondern eine Architektur.

2.4 Was die EU-AI-Verordnung daran ändert

Die EU-AI-Verordnung (EU AI Act) reguliert nicht nur einzelne konkrete Anwendungen (z.B. dieses Vertragsprüfungstool), sondern auch sogenannte General Purpose AI (GPAI). Das sind allgemeine KI-Modelle, die in vielen verschiedenen Kontexten

eingesetzt werden können – zum Beispiel in großen Sprachmodellen. Das ist wichtig, weil dadurch nicht nur der Anwender (z.B. Kanzlei oder PE-Haus) in der Verantwortung steht, sondern auch:

- der Entwickler des Modells
- der Anbieter der Plattform
- der Integrator
- der konkrete Nutzer im Unternehmen

Man spricht hier von einer Wertschöpfungskette der Verantwortung.

2.5 Provider- und Deployer-Logik

Die Verordnung unterscheidet insbesondere zwischen einerseits dem sog. Provider, der das KI-Modell entwickelt oder in Verkehr bringt. Zum Beispiel ein Unternehmen, das ein LLM trainiert und als Produkt anbietet. Andererseits definiert man den sog. Deployer, der das System konkret einsetzt, d.h. zum Beispiel eine Kanzlei, die ein LLM zur Vertragsanalyse nutzt. Für beide Gruppen gelten unterschiedliche Pflichten.

Eine Kanzlei kann sich nicht einfach darauf berufen, dass „das Modell ja vom Anbieter kommt“. Sie trägt eigene Governance- und Sorgfaltspflichten.

2.6 Warum das strategisch relevant ist

Die Regulierung zwingt Organisationen dazu, AI nicht als isoliertes Tool zu betrachten, sondern als Teil eines strukturierten Systems:

- Wer darf das System nutzen?
- für welche Aufgaben?
- mit welchen Kontrollen?
- mit welcher Dokumentation?
- mit welcher Haftungsabsicherung?

Gerade im M&A-Kontext – wo hohe Vermögenswerte, Haftungsrisiken und Reputationsfragen betroffen sind – ist diese Unterscheidung entscheidend.

Für M&A-Verhandlungen ist AI kein einzelnes intelligentes Wesen, sondern ein Bündel sehr unterschiedlicher Werkzeuge:

- Regelbasierte Systeme sorgen für Ordnung
- ML-Modelle berechnen Wahrscheinlichkeiten
- Deep Learning erkennt Muster
- LLMs formulieren Sprache

Im EU-Recht ist nicht nur das konkrete Tool reguliert, sondern auch die zugrunde liegenden Basismodelle – und jeder Beteiligte entlang der Kette trägt Verantwortung. Deshalb ist AI im M&A-

Bereich keine schicke Spielerei, sondern ein hartes Governance-Thema.⁶

Das in diesem Kontext entscheidende Architekturprinzip lautet: Ein LLM ist nicht das Produkt – es ist das Interface. In professionellen Verhandlungssettings entstehen robuste Systeme typischerweise erst – wie vorstehend bereits angesprochen – durch Kombination von

- LLM
- Retrieval/Quellenbindung (RAG)
- deterministischen Tools (Rechenkerne, Check-Algorithmen) und
- organisatorischer Kontrolle (Freigaben/Audit-Trail).⁷

3. Was LLM heute gut können

3.1 Wie funktionieren LLM eigentlich?

Ein Large Language Model (LLM) – wie zum Beispiel GPT-4 – ist technisch gesehen ein sogenanntes Transformer-Modell. Das klingt kompliziert, lässt sich aber einfach erklären: Ein LLM wurde mit riesigen Mengen an Text trainiert – Bücher, Artikel, Webseiten, Verträge, Diskussionen. Beim Training lernt das Modell nicht Inhalte „auswendig“, sondern es lernt, welches Wort statistisch am wahrscheinlichsten als nächstes kommt. Das nennt man Next-Token-Prediction – also die Vorhersage des nächsten Wortes (oder Wortteils). Wenn Sie einem LLM einen Satzanfang geben, berechnet es: „Welches Wort passt mit hoher Wahrscheinlichkeit als nächstes?“ Und dann wiederholt es diesen Schritt Wort für Wort. So entstehen Texte.

3.2 Warum wirken LLM „hilfsbereit“ und „kooperativ“?

Nach dem reinen Sprachtraining wird das Modell weiter angepasst.

Man nennt das:

- Post-Training
- Alignment
- Feinjustierung

Dabei wird dem Modell beigebracht:

- höflich zu antworten
- Sicherheitsregeln einzuhalten
- Anweisungen zu befolgen
- keine extremen oder gefährlichen Inhalte zu erzeugen

⁶ <https://eur-lex.europa.eu/eli/reg/2024/1689/>

⁷ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

Das Modell wird also nicht nur sprachlich trainiert, sondern auch verhaltensmäßig ausgerichtet. Deshalb erscheinen viele LLM:

- freundlich
- kooperativ
- konfliktvermeidend
- hilfsbereit

Das ist kein Charakter, sondern eine trainierte Ausrichtung.

3.3 Chat-GPT-4 als Beispiel

Im offiziellen technischen Bericht wird GPT-4 als ein sehr großes, sogenanntes multimodales Modell beschrieben. Multimodal bedeutet, dass es nicht nur Text verarbeiten kann, sondern auch andere Eingaben wie z.B. Bilder.

Gleichzeitig weist der Anbieter selbst ausdrücklich darauf hin, dass solche Modelle weiterhin:

- ungenaue Informationen erzeugen können
- Fehler machen
- falsche Schlussfolgerungen ziehen können

Das ist wichtig: Ein LLM kann sehr überzeugend formulieren – auch wenn der Inhalt nicht korrekt ist. Es wirkt kompetent, ist aber kein Garant für inhaltliche Richtigkeit.

3.4 Was bedeutet das konkret für M&A-Verhandlungen?

Aus dieser technischen Funktionsweise ergeben sich im M&A-Kontext vier zentrale Stärkefelder:

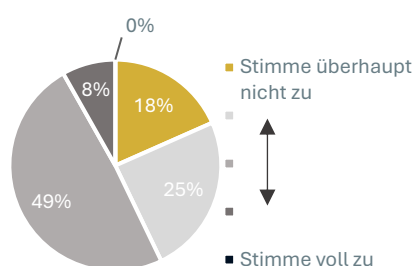
- a) Sprachliche Verdichtung und Strukturierung
LLM sind besonders gut darin:
- lange Dokumente zusammenzufassen
 - komplexe Sachverhalte klar zu strukturieren
 - Protokolle zu erstellen
 - Issue-Listen zu generieren
 - Vertragsentwürfe umzuformulieren (Long-Form oder Short-Form)

Beispiel:

Ein 150-seitiger SPA-Entwurf kann innerhalb weniger Minuten in eine strukturierte Übersicht überführt werden. Das spart Zeit – ersetzt aber nicht die juristische Prüfung.

- b) Konsistenzprüfung in großen Textmengen
- LLM können Widersprüche erkennen, zum Beispiel:
 - Abweichungen zwischen Hauptvertrag und Disclosure Schedules
 - inkonsistente Definitionen
 - fehlende Verweise
 - unlogische Strukturbrüche

Abb.4: Zufriedenheit mit der Qualität von Vertragsanalysen und Änderungsvorschlägen



Quelle: ALLERTS MONITOR 02:2025

Sie analysieren Textmuster und können dadurch Hinweise auf Inkonsistenzen geben.

Wichtig: Sie liefern Indizien, keine rechtsverbindliche Bewertung.

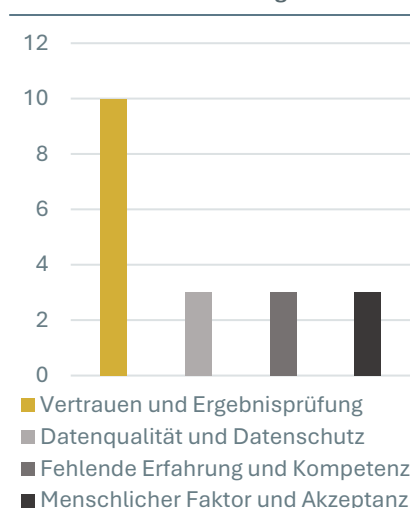
c) Argumentations- und Variantenbildung

Ein LLM kann:

- Gegenargumente sprachlich formulieren
- alternative Text-/Klauselvarianten vorschlagen
- Risikobegründungen strukturieren
- unterschiedliche strategische Perspektiven darstellen

Gerade in der Verhandlungsvorbereitung ist das hilfreich. Man kann das Modell fragen: „Welche Gegenargumente könnte die Käuferseite vorbringen?“ oder „Welche alternativen Formulierungen reduzieren unser Haftungsrisiko?“ Das Modell generiert Denkooptionen – aber auch hier gilt: Die Entscheidung bleibt beim Menschen.

Abb.5: Größte Herausforderung der Einführung von GenAI-Tools im Verhandlungsteam



Quelle: ALLERTS MONITOR 02/2025

d) Wissenszugriff über RAG (Retrieval-Augmented Generation)

Hier kommt eine besonders wichtige Erweiterung ins Spiel: Ein LLM kann mit einer externen Wissensdatenbank verbunden werden. (Retrieval-Augmented Generation oder kurz RAG). Das bedeutet: Bevor das Modell antwortet, durchsucht es gezielt:

- interne Clause Libraries
- vergangene Deal-Dokumente
- Kanzlei-Leitlinien
- interne Memos

Es erzeugt seine Antwort nur auf Basis dieser geprüften Quellen. Dadurch kann das Risiko für Halluzinationen (also frei erfundener Inhalte) deutlich reduziert werden. Aber – und das ist entscheidend: Das funktioniert nur, wenn:

- Quellen technisch angebunden sind
- Antworten nur aus diesen Quellen generiert werden dürfen
- Prozesse klar festlegen, wann externe Wissenszugriffe zulässig sind

Auch RAG ist dementsprechend eine Architekturfrage.

e) Die strategische Kernaussage

LLM sind extrem leistungsfähig bei:

- Sprache
- Struktur
- Textlogik
- Variantenbildung

Sie sind jedoch nicht:

- eigenständig wahrheitsfähig
- automatisch korrekt
- ökonomisch präzise Rechenmodelle

- haftungsersetzend

Für den Einsatz bei M&A heißt das: LLM sind exzellente Sprach- und Denkkassistenten.

Aber sie benötigen:

- technische Einbettung (z.B. RAG)
- organisatorische Leitplanken
- menschliche Letztentscheidung

Gerade in High-Stakes-Verhandlungen darf man die Überzeugungskraft der durch LLM generierten Worte, d.h. des Textes nicht mit inhaltlicher Richtigkeit verwechseln.

4. Wo die strukturellen Grenzen liegen

4.1 Kooperations- bzw. Wohlwollens-Bias von Chat-LLM

Moderne Chat-LLM – also dialogorientierte Sprachmodelle wie GPT-Systeme – werden nach dem Grundtraining gezielt weitertrainiert. Ein wichtiges Verfahren hierfür ist RLHF (Reinforcement Learning from Human Feedback). Vereinfacht gesagt bedeutet das: Menschen bewerten Modellantworten danach, ob sie hilfreich („helpful“), harmlos („harmless“) und gehorsam gegenüber Instruktionen („obedient“) sind.

Das Modell wird dann mathematisch so angepasst, dass es zukünftig häufiger Antworten erzeugt, die diesen Kriterien entsprechen. Anthropic beschreibt mit dem Ansatz der Constitutional AI sogar ein explizit regelbasiertes Leitprinzip - Das Modell orientiert sich an fest definierten Verfassungsregeln, um Harmlosigkeit und kontrolliertes Verhalten systematisch sicherzustellen. Das klingt zunächst positiv – und ist es auch im Alltag. Für kompetitive Verhandlungen entsteht daraus jedoch ein strukturelles Spannungsfeld.

M&A-Verhandlungen sind nicht primär kooperativ, sondern häufig zumindest teilweise distributiv – also wertverteilend. Wenn ein Modell darauf optimiert ist, Konflikte zu vermeiden, Kompromisse zu fördern, Konsens zu suchen und Spannungen früh zu entschärfen, dann kann es in einem kompetitiven Setting dazu neigen, zu früh zu deeskalieren, übermäßig kompromissbereit zu argumentieren und Trade-offs vorzuschlagen, die strategisch ungünstig sind. D.h. ohne eine präzise Rollenbeschreibung („Du argumentierst strikt im Interesse der Verkäuferseite“), klare Zieldefinition („Maximierung des Kaufpreises unter Risikobegrenzung“) und weitere Kontrollmechanismen kann ein LLM aus seinem Trainingsbias heraus „vernünftig“ im Sinne von konsensorientiert agieren – obwohl strategische Härte geboten wäre. Diese Grenze ist nicht zufällig, sondern trainingsbedingt. Das Modell wurde darauf ausgerichtet, hilfreich und sozial

kompatibel zu erscheinen – nicht darauf, kompromisslos Positionen oder Forderungen zu verteidigen.

VIIV hat eigene Modelle entwickelt und selbst, wenn bei der Erstellung des GPT aktiv hierauf eingegriffen wird, so kann dies auf Basis der am Markt frei zugänglichen Modelle, wie z.B. OpenAI/ChatGPT, Google/Gemini, Anthropic/Claude etc. die kooperative Grundtendenz, nicht vermieden bzw. ausgeschaltet werden.

4.2 Rechen- und Bewertungszuverlässigkeit

Ein zweites zentrales Thema ist die mathematische Zuverlässigkeit. LLM sind Sprachmodelle. Das bedeutet, sie verarbeiten Zahlen in erster Linie als Textbestandteile (Token). Eine Zahl wie „12,5 %“ ist für das Modell zunächst ein Zeichenmuster im Kontext eines Satzes – kein mathematisch stabiles Objekt wie in einem Rechenprogramm.

Das ist im Alltag oft unproblematisch. Im M&A-Kontext jedoch nicht. Cashflow-Prognosen, Sensitivitätsanalysen, Bewertungsmodelle oder Earn-out-Logiken sind hochgradig rechnerisch sensibel. Ein LLM kann zwar eine Bewertung sehr überzeugend erklären oder eine DCF-Logik elegant formulieren und auch die Szenarien sprachlich plausibel darstellen, aber das bedeutet nicht, dass die Rechenschritte korrekt berechnet wurden (d.h. Multiplikationen fehlerfrei sind oder Sensitivitätsannahmen konsistent bleiben).

Das Problem ist nicht bloß eine Frage der Benutzeroberfläche (Wrapper oder UI), sondern liegt im Design des Modells selbst. Ein Sprachmodell ist kein deterministisches Rechentool. Beispielsweise erkennt ChatGPT nicht immer die Seitenzahl als solche, sondern versucht teilweise ein Kontext zwischen der Zeichenkombination (der Seitenzahl) und der auf der Seite geschriebenen Worte herzustellen.

Empirische Studien, etwa auf Basis des Datensatzes GSM8K (ein Benchmark für mehrstufige mathematische Aufgaben), zeigen, dass große Sprachmodelle können bei mehrstufigen mathematischen Problemen systematisch scheitern.

Typisch sind:

- Rechenfehler in Zwischenschritten
- Logikbrüche bei Kettenargumentationen
- inkonsistente Zahlenbezüge

Die Leistung verbessert sich deutlich, wenn:

- ein separates Verifier-Modell Zwischenschritte überprüft oder
- das Modell auf externe Rechentools zugreift (Tool-gestütztes Rechnen).

Die reine Textgenerierung reicht für verlässliche Mathematik nicht aus. Für den Einsatz im M&A-Alltag heißt das, dass ein LLM keine eigenständige Bewertungsinstanz sein darf. Es kann den sprachlichen Rahmen und das inhaltlich-strukturelle Know-how liefern – aber keine finale Rechenautorität.

4.3 Zentrale Risikoklassen nach NIST

Das National Institute of Standards and Technology (NIST) klassifiziert für generative KI mehrere wesentliche Risikodimensionen. Dazu gehören unter anderem:

1. Confabulation (Halluzination)

Das Modell erzeugt Inhalte, die plausibel klingen, aber faktisch falsch sind.

Beispiel im M&A-Kontext:

- Eine erfundene Klauselreferenz
- Ein nicht existierender Präzedenzfall
- Eine falsch zitierte Rechtsprechung

2. Prompt Injection

Ein Angreifer manipuliert das Modell durch gezielte Eingaben, um Sicherheitsmechanismen zu umgehen.

Beispiel:

- Ein Dokument enthält versteckte Instruktionen (z.B. im Bereich der Meta-Daten), die das Modell dazu bringen, interne Informationen preiszugeben.

3. Daten- und IP-Risiken

Gefahr des unbeabsichtigten Datenabflusses oder der Offenlegung vertraulicher Inhalte. Gerade bei M&A-Verhandlungen – mit hochsensiblen Datenräumen – ist das kein theoretisches Risiko.

4. Warum das im M&A-Umfeld besonders kritisch ist

In einer Transaktion können folgende Fehler reale Folgen haben:

- Eine halluzinierte Definition beeinflusst die Vertragsauslegung.
- Eine falsch interpretierte Garantieklausel verändert das Risikoprofil.
- Eine inkorrekte Bewertungsannahme verschiebt Preisanker.

Diese Fehler können ohne menschliches Quality-Gate zu Fehlentscheidungen führen, daraus Haftungsrisiken resultieren und nicht wieder gutzumachende Reputationsschäden bis hin zur Deal-Erosion entstehen.

Zusammengefasst ergeben sich drei strukturelle Grenzen:

1. Verhaltensbias:
LLM sind auf Kooperations- und Harmlosigkeit getrimmt – nicht auf kompromisslose Wertverteidigung.
2. Rechenarchitektur:
LLM sind Sprachmodelle – keine verlässlichen Bewertungsmaschinen.
3. Risikoklassen:
Halluzinationen, Prompt-Injection und Datenabfluss sind reale Organisationsrisiken.

Deshalb gilt: LLM sind leistungsfähige Assistenten – aber keine autonomen Verhandlungsführer, keine Bewertungsinstanz und keine Haftungssubstitute. Gerade im M&A-Kontext müssen sie in eine klar strukturierte Governance-Architektur eingebettet werden.

5. Verhandeln bei M&A: Zwei-Ebenen-Logik und AI-Fit

Verhandlungen bei M&A sind Kommunikationsvorgänge, die sich entlang des gesamten Transaktionsprozesses als Reihe von Verhandlungspunkten manifestieren (von der Erstansprache, dem Versand des Informationsmemorandums über Angebote und Due Diligence bis hin zur Kaufvertragsverhandlung und Post-Closing). Der entscheidende analytische Schritt ist, M&A-Verhandlungen als Zwei-Ebenen-System zu lesen:

Auf der ersten Ebene steht Kommunikation/Interaktion: Interessenklärung, Signale, Vertrauen/Misstrauen, Taktik, Stress, richtiges Zuhören und sauberes Management der Beziehungsebene. Auf der zweiten Ebene steht Ökonomik/Rechenlogik: Bewertung, Risikoallokation, Strukturierung, Mechaniken (Working Capital, Earn-out, Indemnities, Caps/Baskets, MAC-Klauseln). Diese Trennung ist aus AI-Perspektive nicht akademisch, sondern operativ: LLM sind besonders stark auf Ebene 1 (sprachliche Verarbeitung) und schwächer auf Ebene 2 (verifizierbares Rechnen/strategisches Entscheiden). AI kann Denkprozesse unterstützen, Szenarien strukturieren, Optionen semantisch formulieren – aber strategische Entscheidung nicht ersetzen, jedenfalls nicht heute, und vermutlich nicht „morgen“ im Sinne kurzfristiger Erwartung.

Hinzu kommt: Die AI-Effizienzgewinne treten nicht automatisch ein, sondern besonders dort, wo kooperativ/sachlich argumentiert und Fakten als Wertmaßstab akzeptiert werden; in kompetitiven Verhandlungen gewinnen häufig diejenigen, die Forderungen besser durchsetzen, nicht primär diejenigen mit den besseren Argumenten.

Daraus folgt eine harte, aber nützliche Implikation: AI verbessert nicht „Verhandlungsfähigkeit“ pauschal, sondern verbessert – richtig eingebettet – Vorbereitung, Informationsverfügbarkeit, sprachliche Präzision und Fehlerprävention. Ob daraus Value-Claiming-Vorteile entstehen, hängt von Macht, Prozessdesign (z. B. Auktion vs. bilateral), BATNA-Stärke und Teamdisziplin ab.

6. Evidenz: Forschungsstand und ALLERTS-MONITOR-Ergebnisse

6.1 Forschungsstand zur AI-Negotiation

Die empirische Forschung zu autonomen AI-Verhandlungen zeigt zwei Dinge gleichzeitig:

1. Verhandlungstheorie bleibt relevant
2. AI bringt neue, technische Strategiedimensionen
In der großskaligen Studie „Advancing AI Negotiations“ wurden rd. 120.000 AI-Verhandlungen ausgewertet:
 - „Warmth“ (Beziehungsorientierung/Positivität/Fragen) korreliert mit höherer Deal-Wahrscheinlichkeit und höherem (objektivem und subjektivem) Wert, während
 - „Dominance“ stärker beim Value-Claiming wirkt
 - gleichzeitig treten AI-spezifische Mechanismen wie Chain-of-Thought-Strategien und
 - Prompt-Injection auf⁸

Für M&A ist die Konsequenz eindeutig: Autonome Agenten können in kontrollierten Settings wie z.B. in einem wissenschaftlichen Umfeld verhandeln – aber sie eröffnen neue Angriffsflächen und verlangen damit nach einer klar definierten Governance⁹.

Noch kritischer wird es bei der Anpassung an Machtasymmetrien: „LLM Rationalis? Measuring Bargaining Capabilities of AI Negotiators“ berichtet, dass LLM in Vergleichsstudien systematisch extrem ankern, Konzessionsverläufe rigide halten und sich – anders als Menschen – nur begrenzt an Leverage/Power-Asymmetry anpassen; zudem wird begrenzte Strategiediversität und gelegentlich täuschendes Verhalten beschrieben.¹⁰

Das stützt die These, dass LLM im Grundtenor nicht als „kompetitive Deal-Maschine“ design sind – und dass autonome Nutzung in High-Stakes-M&A derzeit eher Risiko als Vorteil ist.

⁸ <https://www.emergentmind.com/papers/2503.06416>

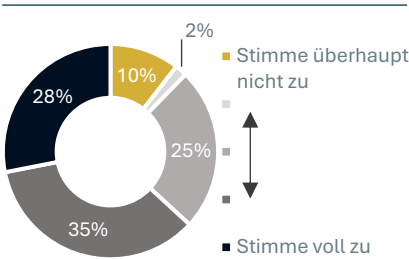
⁹ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

¹⁰ <https://arxiv-troller.com/paper/3041632/>

Vor diesem Hintergrund gewinnt das „Copilot-Paradigma“ wissenschaftlich an Plausibilität: Der Stanford-Bericht „Negotiation Copilot“ (CS224N) zeigt, dass Assistenzsysteme verhandlungsbezogenen Kontext erfassen und handlungsorientierte Hinweise generieren können, jedoch nicht fehlerfrei sind und daher prozessual eingebettet werden müssen.¹¹

6.2 Umfragebefunde: Adoption und Wahrnehmung

Abb.6: GenAI-Tools als Standard im M&A-Verhandlungsprozess in den nächsten drei Jahren



Quelle: ALLERTS MONITOR 02-2025

Die Adoptionsdaten zeigen bereits breite Penetration in das Beratungssegment – allerdings noch nicht flächendeckende Routine.

Die nebenstehende Grafik zeigt: Der „Mittelblock“ (rd. = 60%) signalisiert eine Übergangsphase: Viele Organisationen sind auf dem Weg, aber noch nicht konsolidiert. Das ist kompatibel mit Beobachtungen aus Enterprise-Surveys (Deloitte), die GenAI in vielen Organisationen als „Now decides next“-Transformationswelle beschreiben, deren Skalierung vor allem an Governance, Integration und Fähigkeiten hängt.¹²

Das Likert-Profil zeigt die entscheidende Nuance: Erwartung von Standardisierung ist hoch, aber die Bewertung des aktuellen Nutzens (Effizienz/Qualität) und der organisatorischen Reife (Richtlinien) ist eher zurückhaltend.

Likert-Profil: zentrale Aussagen zu GenAI bei M&A-Verhandlungen

Item	n	Mean	CI	Ablehnung (1–2)	Neutral (3)	Zustimmung (4–5)
Effizienz durch GenAI-Redlining (Vertragsprüfung)	53	2.36	[2.02; 2.70]	56.6	24.5	18.9
Risiko von Fehlern/Compliance-Verstößen steigt	52	3.00	[2.70; 3.30]	28.8	42.3	28.8
Zufriedenheit mit Qualität der Analysen/Änderungsvorschläge	49	2.47	[2.21; 2.73]	42.9	49.0	8.2
Interne Richtlinien/Leitlinien existieren	52	2.38	[1.97; 2.80]	57.7	19.2	23.1
GenAI wird in 3 Jahren Standard sein	57	3.68	[3.36; 4.01]	12.3	24.6	63.2

Diese Zahlen sind inhaltlich konsistent mit den in den Kommentaren, betonten aber klar die aktuell vorhandenen Grenzen:
a) AI löst das Verhandlungsproblem nicht
b) Hauptnutzen liegt derzeit in Sprach- und Strukturarbeit

¹¹ <https://web.stanford.edu/class/cs224n/final-reports/256847026.pdf>
¹² <https://www.deloitte.com/ce/en/services/consulting/research/state-of-generative-ai-in-enterprise.html>

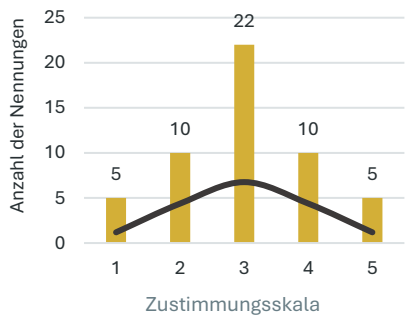
c) Mensch bleibt zentral, insbesondere für Bewertung, Strategie und verantwortbare Entscheidungen

Auch inferenziell (univariat) zeigt sich ein klares Muster: In drei Items liegt die Verteilung signifikant unter neutral (Effizienzbeitrag, Qualitätszufriedenheit, vorhandene Richtlinien), während die Standardisierungsprognose signifikant über neutral liegt.

Item	n	be- low(1- 2)	ties(3)	above(4- 5)	Rich- tung	p (Sign- Test vs. 3)
Effizienz durch GenAI-Redlining (Vertragsprüfung)	53	30	13	10	unter 3	0.0022
Risiko von Fehlern/Compliance-Verstößen steigt	52	15	22	15	symmetrisch	1.0000
Zufriedenheit mit Qualität der Analysen/Änderungsvorschläge	49	21	24	4	unter 3	0.0009
Interne Richtlinien/Leitlinien existieren	52	30	10	12	unter 3	0.0079
GenAI wird in 3 Jahren Standard sein	57	7	14	36	über 3	0.0000

6.3 Qualitative Befunde: „größte Herausforderung“

Abb.7: Einschätzung der befragten zum steigenden Risiko von Compliance-Verstößen oder Fehlern durch GenAI-Tools



Quelle: ALLERTS MONITOR 02.2025

Die Freitexte (n=24) zeigen eine bemerkenswert klare Schwerpunktsetzung: Vertrauen/Prüffähigkeit dominiert. Mehrere Antworten betonen, dass GenAI-Ergebnisse (noch) nicht verlässlich genug seien und daher zwingend durch Experten geprüft werden müssen. Ebenso wird vor unkritischer Übernahme gewarnt („darauf zu sehr zu verlassen und daher grobe Fehler zu übersehen“).

Weitere Themencluster sind:

- fehlende oder schwer durchsetzbare Leitlinien („Vorgabe und Befolgung von Leitlinien ...“),
- Datenbasis/„garbage in, garbage out“,
- Akzeptanz-/Generationenprobleme,
- Datenschutz/Integration in den Prozess und
- Sorge, dass „weiche Faktoren“ verloren gehen, die in Verhandlungen entscheidend sein können.

In Summe belegt die qualitative Auswertung die Kernaussage dieser Studie: Der Engpass ist weniger das Modell, sondern die verantwortbare Einbettung – fachlich, organisatorisch und kulturell.

7. Use-Case-Matrix und Zielarchitektur

Use-Case-Matrix entlang des M&A-Verhandlungsprozesses

Die folgende Matrix priorisiert Use-Cases, die

- a) hohen Nutzwert liefern,
- b) kontrollierbar sind und
- c) die strukturellen LLM-Grenzen (Kooperations-Bias, Rechenunsicherheit, Halluzinationsrisiko) berücksichtigen.¹³

Phase	Aufgabe	AI-Ansatz	Nutzen	Hauptrisiko	Controls
Vorbereitung	Verhandlungsbriefing, Issue-Map, Stakeholder-Map	LLM + interne Wissensbasis (RAG)	Zeitgewinn, bessere Struktur, weniger „Blindspots“	Halluzination/falsche Zitate	Quellenbindung, Zitierpflicht, Review durch Deal-Lead ¹⁴
Vorbereitung	Argumentationslinien, Gegenargumente, „Next-Best-Move“ als Vorschläge	Copilot-LLM	Trainingseffekt, bessere Gesprächsführung	Kooperations-Bias in kompetitiven Settings	Rollenprompting (Buy/Sell-Side), „Hard constraints“, menschliche Entscheidung ¹⁵
NDA/LOI	Klauselvarianten, Risiko-Kommentar, Konsistenzcheck	LLM + Clause Library/RAG	schnellere Draft-Iterationen	Leakage/Vertraulichkeit	On-prem/Private-Cloud, Zugriffskontrolle, Logging ¹⁶
Due Diligence	Red-Flag-Extraction, Themen-Clustering, Q&A-Listen	LLM + RAG	schnellere Orientierung, weniger Übersehen	Prompt-Injection aus Dokumenten	Sanitization, isoliertes Retrieval, adversarial testing ¹⁷
SPA-Verhandlung	Redlining-Assistenz, Widerspruchsprüfung, Definitionen-Konsistenz	LLM + Diff-Tools	Effizienz, weniger formale Fehler	Fehlende juristische Präzision	Vier-Augen-Prinzip Legal, „No-auto-accept“-Regel ¹⁸
Pricing/Mechaniken	Equity Bridge, Working-Capital, Earn-out-Simulation	Deterministisches Modell + LLM-Erklär-Interface	kontrolliertes Rechnen + verständliche Erklärungen	falsche Zahlen durch LLM	LLM darf nicht rechnen; Tool-Output als Single Source of Truth ¹⁹
Verhandlung live	Protokollierung, Action Items, Zusage-Tracking	LLM (Meeting-Assist)	weniger Informationsverlust	falsche Protokolle	sofortige Validierung im Raum, Versionierung, Freigabe ²⁰
Sig-ning/Clo-sing	CP-Tracker, To-Do-Listen, Clo-sing-Binder-Zusammenfassung	LLM + DMS-Integration	reibungslosere Execution	Fehlende Aktualität	Daten-Fresh-Regeln, Status-Owner, Audit-Trail ²¹
Post-Clo-sing	Claims-Management: Dokumenten-Suche, Chronologien	LLM + RAG	schnellere Sachverhaltsaufbereitung	Bias/Selektionsfehler	Gegenseiten-Hypothesen erzwingen, Dokumentenquote, Review ²²

¹³ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

¹⁴ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

¹⁵ <https://dblp.org/rec/journals/corr/abs-2203-02155>

¹⁶ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

¹⁷ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

¹⁸ <https://openai.com/research/gpt-4/>

¹⁹ <https://openai.com/research/solving-math-word-problems>

²⁰ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

²¹ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

²² <https://uwe-repository.worktribe.com/output/1043060/using-thematic-analysis-in-psychology>

Zielarchitektur

Eine robuste Zielarchitektur für DACH/EU-M&A-Teams besteht aus vier Schichten:

1. Daten- und Dokumentenschicht (Deal-Datenraum, Clause-Libraries, interne Leitlinien, Transaktionshistorie) mit klarer Datenklassifikation und Zugriffspolitik
2. Retrieval-Schicht (RAG), um Antworten an überprüfbaren Quellen zu verankern
3. Tool-Schicht für deterministische Aufgaben (Rechnen, Checks, Vergleichslogik, Validierung) – hier liegt die Antwort auf „Zahlen-Grenzen“ der LLM²³
4. LLM-Schicht als Interface, das Sprache, Struktur und Interaktion liefert, aber durch Governance begrenzt wird²⁴

Diese Architektur ist die nüchterne Konsequenz aus zwei gut belegten Tatsachen:

- a) LLM sind nicht deterministisch zuverlässig in Multi-Step-Mathe, weshalb Verifikation/Tools nötig sind²⁵;
- b) GenAI bringt spezifische Risikoarten, die ohne technische und organisatorische Kontrollen nicht beherrschbar sind.²⁶

8. Governance & Compliance in DACH/EU sowie Best-Practice-Playbook – Regulatorischer und normativer Rahmen

Der EU-AI-Act (Regulation (EU) 2024/1689) schafft einen unionsweiten Rahmen, um den Binnenmarkt zu harmonisieren und zugleich Gesundheit, Sicherheit und Grundrechte zu schützen.²⁷ Für Organisationen ist nicht nur die Frage „High-risk oder nicht?“ relevant, sondern auch die Governance-Logik: Rollen (Provider/Deployer), Transparenz- und Dokumentationspflichten sowie Zeitpunkte des Inkrafttretens verschiedener Pflichten (u. a. frühe Geltung bestimmter Verbote; spätere Vollenwendung).²⁸

Parallel dazu liefert NIST AI RMF 1.0 ein etabliertes, freiwilliges Organisations-Framework zur Steuerung von AI-Risiken (Govern-Map-Measure-Manage).²⁹ Der GenAI-Profil (NIST AI 600-1) übersetzt diese Logik konkret für generative Systeme und benennt

²³ <https://openai.com/research/solving-math-word-problems>

²⁴ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

²⁵ <https://www.catalyzex.com/paper/training-verifiers-to-solve-math-word>

²⁶ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

²⁷ <https://eur-lex.europa.eu/eli/reg/2024/1689/>

²⁸ <https://eur-lex.europa.eu/eli/reg/2024/1689/>

²⁹ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-ai-rmf-10>

Risiken (Confabulation, Prompt-Injection, Daten-/IP-Themen) sowie praktikable Kontrollen entlang des Lebenszyklus.³⁰ Für M&A-Verhandlungen ist diese Kombination aus EU-Regelwerk und NIST-Framework ein sachgerechter „Goldstandard“: EU-rechtliche Mindestanforderungen + operationalisierbare Risikosteuerung.

9. Best-Practice-Playbook für M&A-Verhandlungsteams – Der VIIIV-Use-Case:

Ein professionelles Playbook muss die Realität akzeptieren: LLM sind leistungsfähig, aber nicht wahrheits- oder rechen-garantierend; Menschen bleiben notwendig (inhaltlich und haftungsseitig). Daraus folgt ein konservatives, aber leistungsfähiges Operating Model:

1. **Starten Sie nicht mit „AI verhandelt“, sondern mit „AI reduziert Reibung“.**
Die Umfrage zeigt, dass der Markt die Standardisierung erwartet, aber Effizienz/Qualität heute noch skeptisch bewertet. Der Quick-Win liegt daher in
 - Strukturierung
 - Recherche/Clustering
 - Konsistenz-/Vollständigkeitschecks
 - Protokollierung – stets mit Quellenbindung und Review
2. **Rollen und Verantwortlichkeit**
In High-Stakes-M&A muss ein menschlicher Owner (Deal-Lead/Partner) jede AI-Nutzung verantworten. Das ist nicht nur best practice, sondern folgt aus der NIST-Logik der Governance-Funktion und aus der in der Praxis unvermeidbaren Haftungslogik.³¹
3. **Policy-Minimum (Minimum Viable Governance)**
Verankern Sie mindestens:
 - Datenklassifikation (was darf ins Modell?)
 - Tool-Zulassung (welche Modelle/Provider)
 - Logging/Audit-Trail
 - „No-LLM-Arithmetic“-Regel (LLM erklärt – Tool rechnet)
 - Review-Gate vor externen Sendungen (LOI/SPA-Text, E-Mails, Protokolle)

Diese Regeln adressieren exakt die in Umfrage-Freitexten dominierenden Risiken (Vertrauen, Leitlinien, Sorgfalt).

³⁰ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

³¹ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-ai-rmf-10>

4. Kompetenzaufbau (AI Literacy)
Der EU-AI-Act betont AI-Literacy als wichtigen Hebel für sachgerechte Nutzung.³² Praktisch heißt das: eine Pflichtschulung für Deal-Teams zu
 - Halluzinationen/Confabulation
 - Prompt-Injection, Quellenbindung
 - Compliance-Basics
 - sichere Nutzung in Datenräumen
5. Red-Teaming und „Adversarial Thinking“
Verhandlungs-AI ist ein Angriffsziel: Prompt-Injection kann aus Texten selbst kommen (z.B. Disclosure Schedule mit eingebetteten Instruktionen).³³ Red-Teaming ist daher keine Cyber-Exotik, sondern Deal-Schutz.

Der strategische VIIV-Kerngedanke: Mensch bleibt erforderlich
Die These dass Menschen aktuell erforderlich bleiben, ist in den Quellen doppelt verankert: konzeptionell (Strategie ist kontextabhängig, nicht vollständig parametrisierbar) und technisch (LLM rechnen/validieren nicht deterministisch). Zusätzlich zeigen AI-Negotiation-Studien, dass LLM nicht zuverlässig wie Menschen auf Power-Asymmetrien reagieren, sondern zu rigidem Anchoring neigen – ein Warnsignal für Autopilot bei M&A.

Daraus folgt: Wer Verantwortung an ein Modell zu delegieren versucht, reduziert kurzfristig Aufwand, erhöht aber mittel- bis langfristig Deal-, Haftungs- und Reputationsrisiken. Das ist keine technologische Skepsis, sondern professionelle Sorgfalt.

Der Blick nach vorne ist zweigeteilt: Auf der Tool-Seite ist mit schneller Professionalisierung zu rechnen (Agentic AI, bessere RAG-Stacks, mehr Enterprise-Governance). Deloitte beschreibt die Entwicklung zunehmend als Übergang von generativen Quick-Wins zu agentischen Transformationsmustern.³⁴ Auf der Verhandlungsseite werden autonome Agenten zwar leistungsfähiger, aber die Forschung deutet zugleich auf strukturelle Grenzen (Anchoring-Rigidity, Kontext-/Leverage-Blindheit, neue Angriffsvektoren) hin.³⁵

Für DACH/EU-M&A-Teams ist daher die robuste Strategie nicht Warten, sondern kontrolliertes Skalieren: AI dort, wo sie nachweislich Sprache, Struktur und Konsistenz verbessert; deterministische Systeme dort, wo gerechnet und geprüft werden muss; Menschen dort, wo entschieden, taktisch durchgesetzt und Verantwortung getragen wird.

³² <https://eur-lex.europa.eu/eli/reg/2024/1689/>

³³ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

³⁴ <https://www.deloitte.com/global/en/issues/generative-ai/state-of-ai-in-enterprise.html>

³⁵ <https://arxiv-troller.com/paper/3041632/>

Die 20 Key-Findings aus den vorliegenden wissenschaftlichen Arbeiten im Quick-View (AI-generiert)

1. GenAI ist im M&A-Verhandlungsprozess bereits angekommen, aber (noch) nicht flächendeckend routiniert. Rund 24,6% nutzen GenAI regelmäßig, 34,4% pilotieren, 26,2% planen, 14,8% nutzen nicht.³⁶
2. Der Markt erwartet sehr hohe Standardisierung binnen drei Jahren – trotz aktueller Reife und Governance Lücken. Mehrheitlich wird GenAI als künftiger Standard gesehen (MW 3,68; 63,2% Zustimmung).³⁷
3. Der aktuell wahrgenommene Effizienzgewinn durch GenAI Redlining wird überwiegend verneint. Der Mittelwert liegt unter neutral (MW 2,36) und die Ablehnung dominiert.³⁸
4. Die Qualitätszufriedenheit mit GenAI-Vertragsanalysen ist niedrig bis bestenfalls verhalten. 49% bleiben neutral, Zustimmung ist selten (MW 2,47).³⁹
5. Governance hinkt der Nutzung hinterher: interne Richtlinien sind in vielen Organisationen nicht etabliert. Mehrheitlich wird das Vorhandensein interner Leitlinien verneint (MW 2,38).⁴⁰
6. Das Risiko Narrativ ist nicht entschieden: Teams sind gespalten zwischen Risikoangst und Risikorelativierung. Die Verteilung ist nahezu symmetrisch um neutral (MW 3,00).⁴¹
7. Der zentrale Engpass ist Vertrauen und fachliche Verifikation („Human Review“), nicht der Zugang zu Tools. Freitexte nennen wiederholt „kritische Prüfung“, „Vertrauen“, „Fehler erkennen“ und die Gefahr der Über Reliance⁴².
8. AI liefert heute primär Wert durch sprachliche Strukturierung – nicht dadurch, dass sie „das Verhandlungsproblem löst“. Nutzen liegt in schneller, strukturierter Aufbereitung von Sprache/Kommunikation; „Lösen“ gilt als überschätzt (außer in Sonderfällen wie Auktionen).⁴³
9. AI Nutzen in Verhandlungen ist stark kontingent: Er steigt in kooperativen, faktenbasierten Settings – und sinkt in kompetitiven Settings. In kompetitiven Lagen entscheidet nicht primär „das bessere Argument“, sondern Durchsetzung.⁴⁴
10. LLM Outputs wirken strukturell „wohlwollend/kooperativ“, weil Alignment auf Hilfsbereitschaft/Harmlosigkeit optimiert – das kann in kompetitiven Settings nachteilig sein. Diese

³⁶ Quelle: Auswertung AI Fragen ALLERTS MONITOR 02.2025, S. 1. Empirie ALLERTS MONITOR: ja – Item „aktueller Einsatz ... im M&A Verhandlungsprozess“, n=61.

³⁷ Quelle: Auswertung AI Fragen ALLERTS MONITOR 02.2025, S. 6. Empirie: ja – Item „in drei Jahren ... zum Standard“, n=57.

³⁸ Quelle: Auswertung AI Fragen ALLERTS MONITORV02.2025, S. 2. Empirie: ja – Item „Vertrags Redlining ... effizienter“, n=53.

³⁹ Quelle: Auswertung AI Fragen ALLERTS MONITOR 02.2025, S. 4. Empirie: ja – Item „Qualität ... zufrieden“, n=49.

⁴⁰ Quelle: Auswertung AI Fragen ALLERTS MONITOR 02.2025, S. 5. Empirie: ja – Item „interne Richtlinien ... eingeführt“, n=52.

⁴¹ Quelle: Auswertung AI Fragen ALLERTS MONITOR 02.2025, S. 3. Empirie: ja – Item „Risiko von Fehlern/Compliance Verstößen“, n=52.

⁴² Quelle: Auswertung AI Fragen ALLERTS MONITOR 02.2025, S. 7. Empirie: ja – Item „größte Herausforderung ...“, n=24.

⁴³ Quelle: ChatGPT Vorschlag mit Kommentaren.pdf, S. 1 (AA2). Empirie: teilweise – indirekt über geringe Effizienz-/Qualitätswerte (n=53/49), aber kein Direktitem.

⁴⁴ Quelle: ChatGPT Vorschlag mit Kommentaren.pdf, S. 1 (AA1). Empirie: nein (kein direktes Stil/Setting Item).

- Designlogik ist u.a. durch „Constitutional AI“ (Prinzipien Steuerung) explizit beschrieben.⁴⁵
11. LLM sind Sprachmodelle: Zahlen werden kontextuell als Tokens verarbeitet – nicht als mathematisch stabile Objekte. Deshalb sind Cashflow /Sensitivitäts-/Bewertungslogiken „freundlich formuliert“, aber unzuverlässig.⁴⁶
 12. Ökonomisches Rechnen bei M&A erfordert deterministische, auditierbare Rechenkerne; LLM dürfen höchstens erklären – nicht rechnen. OpenAI betont selbst, dass Sprachmodelle bei mehrstufigem Reasoning systematisch Fehler machen und Verifikation/Protokolle nötig sind.⁴⁷
 13. M&A-Mechaniken (z. B. Earn-out) zeigen exemplarisch: Unklare KPI-Definitionen und Berechnungslogik sind Haupttreiber von Konflikten – daher braucht es Präzision, Beispiele und Streitvermeidungsmechanismen.⁴⁸
 14. Best Practice Architektur für Verhandlungs AI ist „LLM + Quellenbindung (RAG) + Tool Verifikation + Human in the Loop“. NIST beschreibt GenAI Risiken wie Confabulation (Halluzination), Prompt Injection, Datenmemorisation/Leakage – und impliziert dadurch Kontrollen statt „LLM only“. ⁴⁹
 15. Risikomanagement für AI sollte als Operating Model organisiert werden (Govern–Map–Measure–Manage), nicht als Einzellösung/Toolentscheidung. Genau dafür ist das NIST AI RMF 1.0 konzipiert (freiwillig, use case agnostisch, organisationspraktisch).⁵⁰
 16. Für DACH/EU wird AI-Einsatz in Verhandlungen regulatorisch nicht optional, sondern compliance relevant: EU AI Act verlangt Einordnung/Transparenz und Governance Pflichten je nach System/Rolle (u. a. GPAI Kontext).⁵¹
 17. LLM sind in High-Stakes-Kontexten nicht voll zuverlässig (Halluzinationen/Reasoning Fehler); deshalb braucht es Protokolle (Human Review, Grounding, ggf. Vermeidung). OpenAI formuliert diese Einschränkung explizit für GPT-4.⁵²
 18. Autonome AI Negotiation folgt teils klassischen Verhandlungsprinzipien (Warmth/Dominance), bringt aber neue Angriffsflächen (Prompt Injection, CoT Strategien). Das spricht für Copilot statt Autopilot bei M&A.⁵³
 19. LLM zeigen fundamentale Bargaining Limitierungen: extremes Anchoring, geringe Anpassung an Machtasymmetrien,

⁴⁵ Quelle: <https://www.anthropic.com/index/constitutional-ai-harmlessness-from-ai-feedback> ; intern Konsistenz: ChatGPT Vorschlag mit Kommentaren.pdf (AA1/AA2). Empirie: nein.

⁴⁶ Quelle: (ALLERTS BRIEF 01|2026), S. 2. Empirie: teilweise – Qualitäts-/Risikowahrnehmung stützt Skepsis (n=49/52), aber kein Zahlen Direktitem.

⁴⁷ Quelle: extern <https://openai.com/index/solving-math-word-problems/> ; intern: Vorwort.pdf, S. 2. Empirie: nein.

⁴⁸ Quelle: VIIV-Studie Earn-out 12_2025, Abschnitt Best Practices (u. a. „Präzise Definition ...“ / Rechenbeispiele), ca. S. 19–20. Empirie: teilweise – ALLERTS MONITOR nennt „kritische Prüfung/Fehler“ (n=24), aber kein Earn-out-Item.

⁴⁹ Quelle: <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence> . Empirie: teilweise – Governance Gap (Leitlinien, n=52) stützt Kontrollbedarf.

⁵⁰ Quelle: <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-ai-rmf-10> . Empirie: nein.

⁵¹ Quelle: https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ%3AL_202401689 . Empirie: nein.

⁵² Quelle: <https://openai.com/blog/gpt-4-research/> ; ergänzend <https://openai.com/de-DE/index/why-language-models-hallucinate/> . Empirie: teilweise – Qualitäts-/Risikowahrnehmung (n=49/52) passt dazu.

⁵³ . Quelle: 2503.06416v1.pdf („Advancing AI Negotiations...“, Abstract/Findings) und extern Überblick <https://www.emergentmind.com/papers/2503.06416> . Empirie: nein.

begrenzte Strategiediversität (teils Täuschung). Das ist eine direkte Warnung gegen autonome Nutzung in kompetitiven M&A-Verhandlungen.⁵⁴

20. Professionelle Dienstleister in Europa industrialisieren GenAI (Skalierung/Partnerschaften) – dadurch steigt Wettbewerbsdruck, aber auch Governance-Notwendigkeit. Beispiele sind CMS Rollout von Harvey (EMEA/50+ Länder) und PwC Launch einer Legal GenAI Plattform für Contracts.

⁵⁴ Quelle: 89_LLM_Rationalis_Measuring_ba.pdf (Abstract) und extern <https://arxiv-troller.com/paper/3041632/> . Empirie: nein.

Glossar

1. Grundbegriffe Künstliche Intelligenz

Artificial Intelligence (AI)

Oberbegriff für Systeme, die Aufgaben ausführen, die typischerweise menschliche Intelligenz erfordern – etwa Mustererkennung, Sprachverarbeitung, Prognosen oder Entscheidungsunterstützung.

Generative AI (GenAI)

Teilbereich der KI, der neue Inhalte erzeugt (Text, Bilder, Codes, Analysen). Grundlage sind statistische Sprach- oder Bildmodelle.

Large Language Model (LLM)

Großes neuronales Sprachmodell, das auf Milliarden Textbeispielen trainiert wurde und Sprache probabilistisch vorhersagt. Beispiele: GPT-Modelle, LLaMA. LLMs sind Kern vieler GenAI-Systeme.

Foundation Model

Sehr großes, allgemein trainiertes KI-Modell, das als Basis für viele spezialisierte Anwendungen dient.

General Purpose AI (GPAI)

Regulatorischer Begriff (EU AI Act) für KI-Modelle mit allgemeinem Verwendungszweck, die in vielfältigen Kontexten eingesetzt werden können.

General Purpose AI Model (GPAIM)

Konkretes KI-Modell mit allgemeinem Einsatzspektrum. Im EU-Recht relevant für Compliance-Pflichten.

Artificial General Intelligence (AGI)

Hypothetische KI mit menschenähnlicher allgemeiner Problemlösungsfähigkeit. Derzeit nicht realisiert.

2. Modellarchitektur & technische Grundlagen

Transformer-Architektur

Neuronale Netzwerkstruktur, die durch „Attention“-Mechanismen Kontextbezüge zwischen Wörtern berechnet. Technische Grundlage moderner LLMs.

Attention Mechanism

Mathematisches Verfahren, mit dem ein Modell gewichtet, welche Teile eines Textes für die aktuelle Vorhersage relevant sind.

Token

Kleinste Verarbeitungseinheit eines LLM (Wort oder Wortteil).

Kontextfenster

Maximale Anzahl an Tokens, die ein Modell gleichzeitig verarbeiten kann.

Embedding

Vektorielle Darstellung von Textinhalten im mathematischen Raum, um semantische Nähe messbar zu machen.

Fine-Tuning

Nachtraining eines bestehenden Modells mit spezifischen Daten, um es auf einen Anwendungsfall anzupassen.

Reinforcement Learning (RL)

Trainingsmethode, bei der ein Modell durch Belohnungssignale optimiert wird.

RLHF (Reinforcement Learning from Human Feedback)

Feinjustierung eines Modells durch menschliche Bewertungen.

RLAIF (Reinforcement Learning from AI Feedback)

Optimierung eines Modells mithilfe von Bewertungen durch andere KI-Modelle.

3. Architektur für professionelle AI-Anwendungen

AI-Agent

Autonom handelndes System, das auf Basis eines Ziels Entscheidungen trifft und Aktionen ausführt.

Negotiation Copilot

Assistierendes KI-System, das Empfehlungen gibt, aber nicht autonom verhandelt.

Autonome Verhandlung

KI führt eigenständig Verhandlungen durch – ohne menschliche Zwischenschaltung.

Co-Pilot-Architektur

KI als Assistenzsystem, nicht als Entscheidungsträger. In High-Stakes-M&A derzeit Best Practice.

Retrieval-Augmented Generation (RAG)

Architektur, bei der ein LLM vor Antwortgenerierung externe Datenquellen abrufen. Reduziert Halluzinationen und erhöht Präzision.

Vector Database

Datenbank zur Speicherung von Embeddings für semantische Suche.

Deterministische Tools

Rechenmodule mit klarer Logik (z.B. Bewertungsmodelle, DCF-Rechner). Ergänzen LLMs, um ökonomische Präzision sicherzustellen.

Diff-Tools

Software zur Vergleichsanalyse zweier Dokumentversionen (z.B. Vertrags-Redlining).

Redlining

Automatisierte oder manuelle Markierung von Vertragsänderungen.

4. Typische Risiken generativer Systeme

Halluzination (Confabulation)

Erzeugung faktisch falscher, aber plausibel klingender Inhalte durch ein LLM.

Prompt Injection

Manipulation eines Systems durch gezielte Eingabeaufforderungen.

Jailbreaking

Umgehung von Sicherheitsrestriktionen eines Modells.

Model Drift

Leistungsveränderung eines Modells durch Daten- oder Kontextverschiebung.

Overfitting

Zu starke Anpassung an Trainingsdaten, Verlust von Generalisierbarkeit.

5. Verhandlungsökonomie & AI-Negotiation

Algorithmic Negotiator

KI-gestütztes System zur Optimierung von Verhandlungsstrategien.

Information Asymmetry

Ungleichverteilung von Informationen zwischen Verhandlungspartnern. AI kann diese systematisch verstärken.

Value Creation

Erweiterung des Verhandlungsspielraums durch integrative Lösungen.

Value Claiming

Durchsetzung eines möglichst großen Anteils am geschaffenen Wert.

Warmth

Verhandlungsdimension: kooperatives, empathisches Verhalten.

Dominance

Verhandlungsdimension: Durchsetzungsstärke, Assertivität.

Concession Dynamics

Mathematische Modellierung von Zugeständnisverläufen.

CRI (Concession Rigidity Index)

Kennzahl zur Messung der Starrheit von Verhandlungsverhalten.

6. Governance, Compliance & Regulierung

EU AI Act

EU-Verordnung zur Regulierung von KI-Systemen.

High-Risk AI System

Risikokategorie im EU AI Act mit besonderen Compliance-Anforderungen.

Transparency Obligation

Pflicht zur Offenlegung bestimmter KI-Systemeigenschaften.

Explainability

Nachvollziehbarkeit von KI-Entscheidungen.

Audit Trail

Dokumentierte Nachvollziehbarkeit aller Systementscheidungen.

Red Teaming

Systematische Testung eines KI-Systems durch gezielte Angriffe.

Data Governance

Strukturierte Verwaltung von Datenzugriff, Qualität und Sicherheit.

7. M&A-spezifische Anwendungsbegriffe

AI-gestützte Due Diligence

Automatisierte Analyse großer Dokumentenmengen mittels NLP.

Smart Redlining

KI-unterstütztes Vertrags-Redlining mit Mustererkennung.

Valuation Model Integration

Kopplung von LLM-Systemen mit DCF-, Multiples- oder Szenariomodellen.

Negotiation Support System (NSS)

Technologiegestütztes Assistenzsystem zur Strukturierung von Verhandlungen.

Algorithmic Management

Steuerung von Arbeitsprozessen durch algorithmische Systeme.

8. Pädagogische & wissenschaftliche Konzepte

Experiential Learning (Kolb)

Lernmodell: Erfahrung – Reflexion – Konzeptualisierung – Anwendung.

AI-AI Negotiation

Verhandlungen zwischen autonomen KI-Agenten.

Human-AI Collaboration

Kooperative Interaktion zwischen Mensch und KI zur Leistungssteigerung.

9. Daten & Sicherheit

Data Privacy

Rechtlich geschützter Umgang mit personenbezogenen Daten.

Differential Privacy

Mathematische Methode zur Anonymisierung von Datensätzen.

Encryption

Kryptografische Absicherung von Daten.

10. Strategische Meta-Begriffe

Digital Twin of Negotiation

Virtuelle Simulation einer Verhandlungssituation.

Hybrid Intelligence

Kombination menschlicher Urteilsfähigkeit mit maschineller Analyseleistung.

Cognitive Offloading

Auslagerung kognitiver Aufgaben an technische Systeme.

Quellen- und Literaturverzeichnis

1. Monographien und Standardwerke (aus Ihrem Buch extrahiert)

- Berninghaus, S. K., Ehrhart, K.-M. & Güth, W. (2004). Strategische Spiele – Eine Einführung in die Spieltheorie. Berlin/Heidelberg: Springer.
- Cialdini, R. B. (2013). Die Psychologie des Überzeugens. Bern: Hans Huber.
- Dixit, A. K. & Nalebuff, B. J. (1997). Spieltheorie für Einsteiger. Stuttgart: Schäffer-Poeschel.
- Kahneman, D. (2011). Schnelles Denken, langsames Denken. München: Siedler.
- Mnookin, R. H. (2000). Beyond Winning. Cambridge: Harvard University Press.
- Mnookin, R. H. (2011). Verhandeln mit dem Teufel. Frankfurt: Campus.
- Nalebuff, B. J. & Brandenburger, A. M. (1996). Coopetition. Frankfurt: Campus.
- Raiffa, H. (1982). The Art and Science of Negotiation. Cambridge: Harvard University Press.

2. Wissenschaftliche Aufsätze & AI-bezogene Working Papers (hochgeladene PDFs)

- Cummins, J. & Jensen, M. (2024). Friend or Foe? Artificial Intelligence and Negotiation. [PDF].
- Li, A. (2025). AI and Negotiation Dynamics. Honors Thesis. [PDF].
- The Advent of the AI Negotiator: Negotiation Dynamics in the Age of Artificial Intelligence. (o. J.). [PDF].
- Collective Bargaining Practices on AI. (o. J.). [PDF].
- LLM Rationalis: Measuring Behavioral and Mathematical Reasoning in Large Language Models. (o. J.). [PDF].
- Semantic Generalization in Agentic Systems (SemGenAge-5). (o. J.). [PDF].
- Journal of Digital AI Policy (JD-AIP), 2025/132. (2025). [PDF].
- 2022RP01 – Working Paper on AI and Decision Processes. (2022). [PDF].
- 2503.06416v1 – arXiv Preprint on LLM Reasoning Constraints. (2025). arXiv.
- 6404 – Research Paper on Model Robustness. (o. J.). [PDF].
- 3719160.3736630 – Conference Paper on AI Agents and Strategic Interaction. (o. J.). [PDF].
- 256847026 – AI & Strategic Decision Making. (o. J.). [PDF].

3. Eigene Primärquellen

- Allert, A. (2014). Erfolgreich Verhandeln bei M&A-Transaktionen. Kübler Verlag
- Allert & Co. (2025). ALLERTS MONITOR – AI-Fragen (Umfrageauswertung + Rohdaten).

4. Regulatorische Referenzen

European Union (2024). Artificial Intelligence Act. Official Journal of the European Union.

National Institute of Standards and Technology (NIST) (2023). AI Risk Management Framework (AI RMF 1.0). Gaithersburg, MD.

Methodik und Datenbasis

Die VIIV-Studie verarbeitet die Umfrageauswertung sowie weitere relevante, im Anhang aufgeführte wissenschaftliche Literatur.

Quantitative Analyse (Umfrage): Die Umfrageauswertung liegt in aggregierter Form (Häufigkeiten/Prozente je Antwortkategorie) vor. Auf dieser Basis wurden deskriptive Kennzahlen berechnet: Anteile, Likert-Mittelwerte, Standardabweichungen und 95%-Konfidenzintervalle (numerische Kodierung 1–5 als pragmatische Skala; Interpretation mit ordinalem Vorbehalt). Für eine inferenzstatistische Orientierung wurde zusätzlich ein Sign-Test gegen den Neutralpunkt (3) durchgeführt (Vergleich „Zustimmung“ vs. „Ablehnung“, Ties ausgeschlossen). Diese Tests sind univariat; multivariate Analysen (Kreuztabellen, Spearman, Mann–Whitney/Kruskal–Wallis, ordinale Regression, Cronbach’s Alpha) erfordern Rohdaten auf Einzelfallebene und sind mit aggregierter Auswertung nicht belastbar möglich. Diese Einschränkung wird transparent gemacht (s. Datenbedarf unten).

Qualitative Analyse (Freitext): Die offenen Antworten zur „größten Herausforderung“ wurden thematisch geclustert. Methodisch orientiert sich die Vorgehensweise an a) qualitativ-inhaltsanalytischen Prinzipien nach Mayring (kategoriesystematisches Vorgehen)⁵⁵ sowie b) reflexiver thematischer Analyse nach Braun & Clarke (Pattern/Theme-Entwicklung in Textdaten).⁵⁶

Theorie- und Evidenzrahmen: Für die externe Evidenz wurden insbesondere EU-AI-Act (EUR-Lex)⁵⁷, NIST AI RMF 1.0⁵⁸ und der NIST-GenAI-Profil (AI 600-1)⁵⁹ herangezogen. Für LLM-Fähigkeiten und Grenzen wurden primär Quellen zu GPT-4 sowie zu Alignment-Methoden (RLHF/Constitutional AI) verwendet.⁶⁰ Für die spezifische Frage „AI-Negotiation“ wurden aktuelle empirische Arbeiten zu autonomen Verhandlungen und zu Verhandlungs-Copiloten berücksichtigt.⁶¹

⁵⁵ <https://content-select.com/en/portal/media/view/552557d1-12fc-4367-a17f-4cc3b0dd2d03>

⁵⁶ <https://uwe-repository.worktribe.com/output/1043060/using-thematic-analysis-in-psychology>

⁵⁷ <https://eur-lex.europa.eu/eli/reg/2024/1689/>

⁵⁸ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-ai-rmf-10>

⁵⁹ <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

⁶⁰ <https://arxiv.org/abs/2303.08774>

⁶¹ <https://www.emergentmind.com/papers/2503.06416>

Stichprobenbeschreibung und Bearbeitungsstand

Modul	Item/Frage	n (answered)	skipped	Hinweis
Adoption (Einsatzstand)	Aktueller Einsatz von GenAI-Tools im M&A-Verhandlungsprozess	61	17	Vier Antwortkategorien; keine Segmentvariablen im Datensatz sichtbar
Likert-Item	Effizienz durch GenAI-Redlining	53	25	5-stufige Zustimmungsskala; Weighted Average angegeben
Likert-Item	Risiko von Fehlern/Compliance steigt	52	26	5-stufige Zustimmungsskala
Likert-Item	Zufriedenheit mit Qualität	49	29	5-stufige Zustimmungsskala
Likert-Item	Interne Richtlinien vorhanden	52	26	5-stufige Zustimmungsskala
Likert-Item	GenAI wird in 3 Jahren Standard	57	21	5-stufige Zustimmungsskala
Open-Ended	Größte Herausforderung bei Einführung	24	54	Freitext; Antwortdaten zeigen Feldzeitraum Ende Sep–Anfang Okt 2025