

KI und der Anwalt

Lernen, nach dem zu suchen, was falsch sein könnte

Rechtsanwalt Joseph Weinrauch, Tel Aviv | Juni 2026

Künstliche Intelligenz hat sich leise in der anwaltlichen Praxis eingenistet. Für die meisten Anwälte vollzog sich der Einstieg schrittweise – ein E-Mail-Entwurf hier, eine Dokumentenzusammenfassung dort. Dann Verträge. Dann Rechtsgutachten. Dann Schriftsätze. Bevor wir den Wandel überhaupt wahrnahmen, war KI Teil des Arbeitsablaufs geworden.

In vielerlei Hinsicht ähnelt sie einem außergewöhnlich effizienten Juniorjuristen: Sie entwirft ohne Ermüdung, fasst zusammen ohne Klagen, schlägt Rechtsprechung mit Leichtigkeit vor und ist um drei Uhr morgens erreichbar. Was ist also das Problem mit diesem neuen Assistenten? Es ähnelt sehr dem Problem von damals: Genauigkeit – jedoch mit weitaus größerer Selbstsicherheit.

Als ich ein junger Rechtsanwalt im Büro meines Ausbilders war und Unterlagen für ihn vorbereitete, hatte ich Angst, etwas zu erfinden. Abkürzungen nahm ich vielleicht – aber nur an den Rändern. Der Kern musste stimmen. Die Angst war der Schutzwall: Angst, entdeckt zu werden, Angst vor Peinlichkeit, Angst vor berufsrechtlichen Folgen. Diese Angst war keine Schwäche – sie war ein beruflicher Filter. KI kennt keine solche Angst. Kein Schamgefühl. Kein Zögern. Sie weiß nicht, was sie nicht weiß, und es ist ihr gleichgültig. Sie wird ein Urteil des Obersten Gerichts erfinden – mit Namen, Datum, Aktenzeichen, Leitsatz –, und es Ihnen in makelloser, sicherer Prosa präsentieren, im Ton nicht zu unterscheiden von einem tatsächlich existierenden Zitat. Sie tut dies ohne Zittern. Und dann wartet sie auf Ihre Zustimmung.

Das Problem hat einen Namen

KI-Halluzinationen entstehen, wenn ein System Informationen erzeugt, die völlig glaubwürdig wirken – gut formuliert, strukturiert, plausibel –, aber sachlich falsch, erfunden oder nicht belegt sind. In der anwaltlichen Praxis ist dies keine theoretische Sorge. Es ist eine standesrechtliche.

Im Juni 2024, als mein Kollege Pim Betist und ich gemeinsam in Amsterdam zum Einsatz von KI in der grenzüberschreitenden Rechtsarbeit referierten, formulierten wir dieses Paradox als zentrale Herausforderung:

„Wenn man sie benutzt, gibt es Ärger. Wenn man sie nicht benutzt, wird es doppelt so viel Ärger geben.“

– Pim Betist, Amsterdam Accord, Juni 2024

Zu diesem Zeitpunkt hatte eine Studie der Stanford University gerade ergeben, dass KI bei 51 % der juristischen Recherchanfragen halluziniert – mehr als die Hälfte. Diese Zahl schockierte das

Publikum. Sie war die Zahl, die das Gespräch veränderte. Und doch lagen die Fälle, die Schlagzeilen machen sollten, noch vor uns.

Die israelische Lektion

Am 20. Februar 2025 erließ Israels Oberstes Gericht ein wegweisendes Urteil, das Juristen im ganzen Land nicht ignorieren konnten. Richterin Gila Canfy-Steinitz befasste sich mit einem Antrag in einer Scheidungssache nach Scharia-Recht, in dem der Anwalt der Antragstellerin KI-generierte Schriftsätze eingereicht hatte, die Folgendes enthielten:

OBERSTES GERICHT ISRAEL — FEBRUAR 2025

- 36 zitierte Entscheidungen des Obersten Gerichts – intern widersprüchlich
- Verweise auf nicht existierende Urteile
- Zitate, die in keiner juristischen Datenbank erschienen

Das Gericht stellte fest, dass der Anwalt blind auf eine unbenannte, von Kollegen „empfohlene“ Website vertraut hatte, ohne jede eigenständige Überprüfung. Richterin Canfy-Steinitz schrieb, der KI-halluzinierte Output habe so glaubwürdig gewirkt, dass der Anwalt es nicht für nötig hielt, dessen Richtigkeit zu prüfen. Das Gericht sanktionierte den Anwalt nicht – mit dem Hinweis, dies sei der erste derartige Fall vor ihm –, warnte jedoch ausdrücklich, dass künftige Fälle alle richterlichen Mittel zu spüren bekämen.

Diese Warnung erwies sich fast sofort als prophetisch. Wenige Tage später erreichte das Oberste Gericht ein weiterer Antrag – diesmal zu einem Gesetz, das die Euthanasie streunender Hunde erleichterte –, der wiederum halluzinierte Fälle enthielt. Das Gericht wies ihn rundweg ab. Kurz darauf beantragte ein Schuldner beim Vollstreckungs- und Inkassobüro unter Berufung auf nicht existierende Fälle. Seine entschuldigenden Nachschriften enthielten – bemerkenswert – weitere KI-halluzinierte Zitate. Er wurde mit 1.000 NIS bestraft. Im Mai 2025 verwarf das Bezirksgericht Tel Aviv einen Sammelklageantrag gegen Krankenkassen aus demselben Grund.

Israel war nicht allein. Ein Forscher, der KI-Halluzinationsfälle in Gerichtsschriftsätzen weltweit verfolgt, hat inzwischen über 1.600 dokumentierte Fälle seit 2023 identifiziert, die überwiegende Mehrheit allein im Jahr 2025. In den Vereinigten Staaten wurden Anwälte sanktioniert, suspendiert und mit Bußgeldern belegt. Ein Anwalt aus Alabama wurde wegen der Einreichung von Schriftsätzen mit gefälschten Zitaten sanktioniert – und zitierte unmittelbar im nächsten Satz nach der Korrektur erneut einen nicht existierenden Fall.

Quellen: Library of Congress, Juni 2025; Lexology, März–April 2025; Damien Charlotin, AI Hallucination Cases Database, 2025; Scientific American, 2025.

Was die Forschung heute zeigt

Die Daten zwischen 2024 und 2026 erzählen eine Geschichte bedeutsamer Verbesserungen, die jedoch noch weit hinter professioneller Zuverlässigkeit zurückbleiben.

Die Stanford-Studie von 2024 – „Large Legal Fictions“ – testete allgemeine Sprachmodelle anhand von über 800.000 überprüfbaren Rechtsfragen. Die Halluzinationsraten lagen zwischen 58 % bei GPT-4 und 88 % bei älteren Modellen. Eine zweite Stanford-Studie, veröffentlicht 2025, unter-

suchte dann die premium-spezialisierten Rechtsinstrumente – Lexis+ AI, Westlaw AI-Assisted Research und Thomson Reuters's Ask Practical Law AI – die Werkzeuge, für die Anwälte tatsächlich professionelle Tarife bezahlen:

STANFORD / JOURNAL OF EMPIRICAL LEGAL STUDIES, 2025 — HALLUZINATIONSRATEN

Lexis+ AI	17 %
Westlaw AI-Assisted Research	33 %
GPT-4 (Basiswert)	43 %

Schlussfolgerung: „Legal RAG can reduce hallucinations compared to general-purpose AI, but hallucinations remain substantial, wide-ranging, and potentially insidious.“ — Magesh et al., „Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools,“ Journal of Empirical Legal Studies, 2025.

Um diese Zahlen einzuordnen: Über alle Bereiche hinweg liegt die Gesamthalluzinationsrate für KI selbst bei den leistungsstärksten Modellen bei rund 6,4 % für juristische Informationen – verglichen mit lediglich 0,8 % bei allgemeinen Wissensfragen. Das Recht ist kein Bereich, in dem KI leicht unterdurchschnittlich abschneidet. Sie schneidet dramatisch und spezifisch unterdurchschnittlich ab, weil juristische Argumentation eine kontextuelle Synthese von Autoritäten erfordert, die Sprachmodelle nicht zuverlässig leisten können.

HALLUZINATIONSRATEN NACH FACHGEBIET — LEISTUNGSSTÄRKSTE MODELLE (2024–2025)

Allgemeinwissen	0,8 %
Finanzdaten	2,1 %
Medizin / Gesundheitswesen	4,3 %
Juristische Informationen	6,4 %

Quelle: AI Hallucination Report 2025–2026, AIAboutAI / drainpipe.io compilation.

Die Verbesserung ist real. Die besten aktuellen Modelle halluzinieren weniger als Modelle von vor zwei Jahren – manche um bis zu 64 % weniger. Aber eine Genauigkeitsrate von 85 % bei der juristischen Zitierung ist kein Bestehen. Es ist ein standesrechtliches Risiko.

Eine andere Denkweise

Die tiefere berufliche Herausforderung ist nicht technischer Natur. Sie ist kognitiv – und ermüdender, als es zunächst erscheint.

Den Großteil meiner Laufbahn folgte die anwaltliche Arbeit einem vertrauten Rhythmus. Ich sammelte Expertise durch Tausende von Dokumenten, Hunderte von Transaktionen. Wenn eine neue Sache eintraf, erinnerte ich mich an einen Präzedenzfall, rief ein ähnliches Dokument ab, passte es an. Die Arbeit war kreativ, aber geerdet: Ich arbeitete mit Material, das ich bereits verinnerlicht hatte, und passte es an das Neue an.

KI kehrt dies vollständig um. Sie liefert nahezu sofort einen gut strukturierten Entwurf. Das Dokument sieht gut aus. Die Zitate erscheinen korrekt. Die Argumentation fließt. Und genau hier beginnt die Gefahr.

Die Fähigkeit, die der Anwalt heute benötigt, ist nicht das Abrufen – es ist das Erkennen. Anstatt im Gedächtnis nach dem richtigen Präzedenzfall zu suchen, muss der Anwalt den Output der KI systematisch hinterfragen:

- › Existieren diese zitierten Fälle tatsächlich in der Datenbank?
- › Besagen die Entscheidungen das, was die KI behauptet?
- › Ist der Gesetzgebungsverweis auf den aktuellen Stand des Gesetzes korrekt?
- › Hat die KI ein selektives Bild gezeichnet – gegenteilige Autoritäten ausgelassen?
- › Bezieht sich die Argumentation auf die tatsächlichen Fakten dieses Falls, oder ist sie generisch?

Dies ist nicht dieselbe geistige Anstrengung wie das Entwerfen. Es kommt eher einem Prüfen gleich. Und es hat eine besondere Qualität, die es genuinen Aufwand bereitet: Man sucht nach Fehlern in einem Output, der korrekt aussieht. Die KI markiert ihre Fehler nicht. Sie präsentiert ein erfundenes Zitat mit demselben selbstsicheren Stil wie ein überprüfbares. Der Anwalt muss die Skepsis bereitstellen, die dem Werkzeug vollständig fehlt.

Diese Verlagerung des kognitiven Aufwands ist real, und sie kann ermüdend sein, gerade weil sie dem beruflichen Instinkt entgegenläuft. Erfahrene Anwälte vertrauen ihrer Lektüre eines Dokuments. Sie haben dieses Vertrauen über Jahrzehnte verdient. Bei KI-Output muss dieser Instinkt bewusst ausgesetzt und durch systematische Überprüfung ersetzt werden – jedes Mal, bei jedem Dokument, unabhängig davon, wie glaubwürdig der Entwurf wirkt.

Fehler lassen sich nur erkennen, wenn man über eine solide Grundlage juristischer Kenntnisse verfügt. KI macht juristische Expertise nicht weniger notwendig. Sie macht kontinuierliches Lernen notwendiger denn je.

Sollten Anwälte KI nutzen?

Ja. Mit angemessener Disziplin eingesetzt, spart KI echte Zeit. Ein schnellerer erster Entwurf bedeutet mehr Zeit für die Analyse. Ein zusammengefasster Dokumentensatz bedeutet mehr Zeit für die Strategie. Die Effizienzgewinne sind real, und in vielen Fällen führen sie zu niedrigeren Kosten für Mandanten.

Doch das Wort Disziplin trägt in diesem Satz ein erhebliches Gewicht. Verantwortungsvoller Einsatz bedeutet:

- › Jeden KI-Entwurf als Ausgangspunkt zu behandeln, niemals als fertiges Produkt
- › Jede zitierte Autorität eigenständig anhand offizieller Datenbanken zu überprüfen
- › KI-generierte Schriftsätze niemals einzureichen, ohne sie so sorgfältig zu lesen, als hätte man sie selbst verfasst – denn vor Gericht gelten Sie als deren Verfasser
- › Zeit in die Qualität der Prompts zu investieren: präzise, kontextualisierte Anweisungen erzeugen erheblich bessere Ergebnisse und reduzieren den Überprüfungsaufwand

Die vielleicht treffendste Analogie lautet: KI ist ein außerordentlich effizienter Bibliothekar – einer, der Informationen in Sekunden abrufen kann. Die Verantwortung für das Verstehen, Überprüfen

und Anwenden dieser Informationen liegt vollständig beim Anwalt. Der Bibliothekar kann nicht sanktioniert werden. Sie schon.

Das aktuelle Bild

Zwei Jahre nach Amsterdam hat die Glaubwürdigkeit von KI in der anwaltlichen Praxis zugenommen – messbar, bedeutsam –, jedoch noch nicht genug für den unbeaufsichtigten Einsatz. Die Halluzinationsraten in spezialisierten juristischen Werkzeugen sind von rund 55 % auf 15–20 % gesunken. Das ist Fortschritt. Es ist noch keine Sicherheit.

Die Gerichte haben die Grenze klar gezogen. Rechtsanwaltskammern von Kalifornien über New York bis Tel Aviv haben dieselbe Botschaft ausgegeben: KI verringert die berufliche Verantwortung nicht. Sie schafft eine neue Ebene davon.

Der Anwalt des kommenden Jahrzehnts wird zwei unterschiedliche Kompetenzen beherrschen müssen, die selten nebeneinander bestehen mussten: tiefes juristisches Wissen und rigorose Skepsis gegenüber den Werkzeugen, die behaupten, bei dessen Anwendung zu helfen. Diese Kombination – Expertise und strukturierter Zweifel – ist das, was verantwortungsvolle KI-gestützte Praxis im Juni 2026 bedeutet.

In einer unserer Arbeitssitzungen zu diesem Artikel teilte ich Claude eine Entwurfsrichtung mit. Die Antwort lautete: „Love this direction — personal, authentic, and sharp. Let me complete it in your voice.“
(im Original auf Englisch)

Von einer KI – das ließ mich lächeln.

Quellen: Stanford RegLab / Institute for Human-Centered AI (Januar 2024); Magesh et al., „Hallucination-Free?“, *Journal of Empirical Legal Studies* (2025); Library of Congress Global Legal Monitor (Juni 2025); Lexology (März–April 2025); Damien Charlotin, *AI Hallucination Cases Database* (2025); *AI Hallucination Report 2025–2026* (AllAboutAI); *Scientific American* (Oktober 2025). Amsterdam Accord presentation: Adv. Joseph Weinrauch & Pim Betist, 7. Juni 2024.
Dieser Artikel wurde mit der menschlichen Unterstützung von Frau Preeti Desi und mit Claude AI erstellt.