

# AI Briefing

Februar 2026



**+++ OLG HAMBURG: TEXT UND DATA MINING FÜR KI-TRAINING ZULÄSSIG BEI FEHLERHAFTEM NUTZUNGSVORBEHALT +++ LG FRANKFURT: HAFTUNG VON GOOGLE FÜR "KI-ÜBERSICHT" IST NICHT PER SE AUSGESCHLOSSEN +++ AG MÜNCHEN: ZUM URHEBERRECHTSSCHUTZ FÜR KI-GENERIERTE GRAFIKEN +++ EU-KOMMISSION: UNTERSUCHUNG GEGEN X WEGEN SYSTEMISCHER RISIKEN DURCH GROK EINGELEITET +++ EU-KOMMISSION: ERSTER ENTWURF FÜR KI-TRANSPARENZ-KODEX VERÖFFENTLICHT +++ STANFORD-STUDIE: KI-MODELLE REPRODUZIEREN URHEBERRECHTLICH GESCHÜTZTE BÜCHER FAST VOLLSTÄNDIG +++**

## 1. Rechtsprechung

**+++ OLG HAMBURG: TEXT UND DATA MINING FÜR KI-TRAINING ZULÄSSIG BEI FEHLERHAFTEM NUTZUNGSVORBEHALT +++**

Das Oberlandesgericht Hamburg hat die Berufung eines Fotografen gegen einen gemeinnützigen Verein zurückgewiesen, der ein Bild des Fotografen für KI-Trainingszwecke genutzt hatte. Der Beklagte hatte im Rahmen der Erstellung eines Datensatzes mit 5,85 Milliarden Bild-Text-Paaren die streitgegenständliche Fotografie von einer Bildagentur heruntergeladen, um einen Abgleich zwischen Bild und Bildbeschreibung vorzunehmen. Das Gericht bestätigte, dass sich der Verein auf die Schrankenregelung für Text und Data Mining aus § 44b UrhG berufen könne. Entscheidend sei, dass der auf der Webseite der Bildagentur vorhandene Nutzungsvorbehalt nicht die gesetzlich vorgesehene Form der Maschinenlesbarkeit aufweise, weshalb die Vervielfältigung zulässig sei. Zusätzlich stellt das Gericht fest, dass die Nutzung auch unter die Schrankenregelung für wissenschaftliche Forschung nach § 60d UrhG falle, da bereits die Erstellung des

Datensatzes ein methodisches, auf Erkenntnisgewinn gerichtetes Vorgehen darstelle. Der Umstand, dass auch kommerzielle Anbieter den Datensatz nutzen könnten, ändere daran nichts, da es an einem bestimmenden Einfluss eines privaten Unternehmens fehle. Das Urteil ist nicht rechtskräftig, die Revision wurde zugelassen.

[Zum Urteil des Oberlandesgerichts Hamburg \(v. 10. Dezember 2025, 5 U 104/24\).](#)

### **+++ LG FRANKFURT: HAFTUNG VON GOOGLE FÜR "KI-ÜBERSICHT" IST NICHT PER SE AUSGESCHLOSSEN +++**

Das Landgericht Frankfurt hat entschieden, dass Suchmaschinenbetreiber grundsätzlich für falsche Informationen in KI-generierten Übersichten haften können, im konkreten Fall jedoch keinen Unterlassungsanspruch gegen Google gewährt. Ein Ärzteverband hatte gegen Google geklagt, weil bei Suchen zum Thema "Penisverlängerung" in der vorangestellten "KI-Übersicht" eine medizinisch falsche Aussage erschien. Dadurch sei sein Geschäft beeinträchtigt worden, da es zu weniger Klickzahlen komme und Nutzer wegen der (im vorliegenden Fall aus Sicht des Arztes inhaltlich unzutreffenden) KI-Übersicht seltener auf externe Treffer klickten. Das Gericht sah jedoch keine unbillige Behinderung, da die angegriffene „KI-Übersicht“ im Gesamtkontext aus Sicht der angesprochenen Adressaten nicht als falsch zu bewerten sei. Das Gericht ließ offen, ob Google sich darauf berufen kann, dass es sich bei der KI-Übersicht nur um Informationen Dritter handelt, oder ob sie als eigene Äußerung anzusehen ist. Einen Verstoß gegen den Digital Markets Act lehnte das Gericht ab, da die KI-Übersicht Teil des Suchergebnisses darstelle und nicht auf ein eigenes Produkt von Google verweise (Verbot der Selbstbevorzugung).

[Zum Urteil des Landgerichts Frankfurt \(v. 10. September 2025, 2-06 O 271/25\).](#)

### **+++ AG MÜNCHEN: ZUM URHEBERRECHTSSCHUTZ FÜR KI-GENERIERTE GRAFIKEN +++**

Das Amtsgericht München hat entschieden, dass mittels generativer KI erstellte Grafiken mangels persönlicher geistiger Schöpfung keinen Urheberrechtsschutz genießen. Im vorliegenden Fall ging es um drei verschiedene Logos, die der Kläger durch Text-Prompts unterschiedlicher Komplexität und Detaillierungstiefe generieren ließ. Laut Gericht genügt der bloße Auftrag an die KI („Prompting“) jedoch nicht, um eine schöpferische Prägung durch den Menschen zu begründen, da die eigentliche visuelle Ausgestaltung autonom durch die Software erfolge. Solange der Nutzer nur abstrakte Vorgaben mache und die Umsetzung der KI überlasse, fehle es am erforderlichen menschlichen

Gestaltungsspielraum. Ein urheberrechtlicher Schutz sei aber denkbar infolge menschlichen Eingriffs in die KI-Ergebnisse, der auch nachträglich bzw. sukzessive während des Promptings stattfinden könne und der dazu führe, dass sich im Output auch gerade die Persönlichkeit des Promptenden widerspiegle. Hierfür müsse der menschliche Einfluss den resultierenden Output jedoch hinreichend objektiv und eindeutig identifizierbar prägen. Im Übrigen merkt das Gericht an, dass der Rechtsstreit den Eindruck erwecke, von den Parteien weniger aus wirtschaftlichen Gründen, sondern primär aus wissenschaftlichem Interesse geführt zu werden, um eine gerichtliche Klärung dieser Frage zu provozieren.

[Zum Urteil des AG München \(v. 13. Februar 2026, 142 C 9786/25\).](#)

### **+++ VG HAMBURG: CHATGPT-NUTZUNG ALS TÄUSCHUNGSVERSUCH AUCH OHNE EXPLIZITES VERBOT +++**

Das Verwaltungsgericht Hamburg hat festgestellt, dass die Nutzung von generativer KI (wie ChatGPT) bei schulischen Leistungen als Täuschungshandlung gewertet werden kann, selbst wenn kein ausdrückliches Verbot in der Schulordnung verankert ist. Im konkreten Fall hatte ein Schüler der 9. Klasse für ein englischsprachiges Lesetagebuch („Reading Log“) KI-Unterstützung genutzt, was der Lehrkraft durch stilistische Abweichungen auffiel. Das Gericht bestätigte die Bewertung mit „ungenügend“ und lehnte den Eilantrag des Schülers ab. Zur Begründung führte die Kammer aus, dass generative KI die Eigenständigkeit der Leistungserbringung aufhebe. Bereits die Aufgabenstellung „use your own words“ schließe die Nutzung von Textgeneratoren logisch aus, sodass sich Schüler nicht auf fehlende Detailregelungen zur KI-Nutzung berufen können. Das Gericht stellte zudem klar, dass auch eine nur unterstützende Nutzung (etwa zur sprachlichen Glättung) unzulässig sei, wenn die sprachliche Gestaltung Teil der Bewertungsleistung darstellt.

[Zum Beschluss des VG Hamburg \(v. 15. Dezember 2025, 2 E 8786/25\).](#)

## **2. Behördliche Maßnahmen**

### **+++ EU-KOMMISSION: UNTERSUCHUNG GEGEN X WEGEN SYSTEMISCHER RISIKEN DURCH GROK EINGELEITET +++**

Die Europäische Kommission hat ein förmliches Verfahren gegen X (ehemals Twitter) auf Grundlage des Digital Services Act (DSA) eröffnet. Gegenstand der Untersuchung ist die Prüfung, ob X die mit der Integration des KI-Chatbots „Grok“ verbundenen systemischen Risiken

unzureichend bewertet und gemindert hat. Konkret befürchtet die Behörde, dass Grok zur Erstellung und Verbreitung illegaler Inhalte genutzt werden kann, insbesondere von sexualisierten Deepfakes und Darstellungen, die Kindesmissbrauch ähneln. Zudem wird untersucht, ob die Umstellung auf Grok-basierte Empfehlungssysteme ohne angemessene Risikobewertung erfolgte. Als „Very Large Online Platform“ (VLOP) mit über 45 Millionen Nutzern unterliegt X strengen Auflagen unter dem DSA (Art. 34, 35) und muss vor der Einführung neuer Funktionen zwingend systemische Risiken analysieren und mindern. Die Kommission vermutet, dass diese Prüfung vor dem Start von Grok versäumt wurde.

[Zur Pressemitteilung der EU-Kommission \(v. 26. Januar 2026, IP/26/203\).](#)

## 3. Stellungnahmen / Sonstiges

### **+++ GEMEINSAME STELLUNGNAHME VON EDSA UND EDSB ZUM DIGITALEN OMNIBUS ZU KI +++**

Der Europäische Datenschutzausschuss (EDSA) und der Europäische Datenschutzbeauftragte (EDSB) haben eine gemeinsame Stellungnahme zur Vereinfachung der Umsetzung harmonisierter Regeln für Künstliche Intelligenz im Rahmen des Digitalen Omnibus veröffentlicht. Darin legen sie unter anderem einen Schwerpunkt auf den Ausgleich zwischen Vereinfachung und Datenschutz. So fordern EDSA und EDSB klare Grenzen und hohe Standards für die geplante Erweiterung der Verarbeitung besonderer Kategorien personenbezogener Daten zur Erkennung und Korrektur von Bias in KI-Modellen, damit Missbrauch ausgeschlossen wird. Auch bei der Registrierungs- und Dokumentationspflicht für KI-Systeme warnen sie davor, zentrale Transparenz- und Verantwortlichkeitsmechanismen ersatzlos zu streichen. Die gemeinsame Stellungnahme spricht sich zudem für die Einrichtung von EU-weiten KI-Regulierungs-Sandboxen zur Förderung von Innovation aus, fordert aber mehr rechtliche Klarheit über die Rolle der Datenschutz-Aufsichtsbehörden in diesen Strukturen.

[Zur gemeinsamen Stellungnahme des EDSA und EDSB \(v. 20. Januar 2026, Englisch\).](#)

### **+++ EU-KOMMISSION: ERSTER ENTWURF FÜR KI-TRANSPARENZ-KODEX VERÖFFENTLICHT +++**

Die Europäische Kommission hat den ersten Entwurf des „Code of Practice on Transparency of AI-Generated Content“ vorgelegt, der die Kennzeichnungspflichten nach Art. 50 KI-VO konkretisiert. Das Dokument

ist in zwei Teile gegliedert, die von zwei verschiedenen Arbeitsgruppen erarbeitet wurden. Die „Section I“ richtet sich an Anbieter von KI-Systemen (Art. 50 Abs. 2 KI-VO) und fokussiert sich auf technische Maßnahmen wie Wasserzeichen und Metadaten zur maschinenlesbaren Markierung. Demgegenüber behandelt „Section II“ die Betreiber der Systeme (Art. 50 Abs. 4 KI-VO) und adressiert die sichtbare Kennzeichnung von Deepfakes sowie Texten von öffentlichem Interesse. Der Kodex fordert insgesamt einen „Multi-Layered Approach“, bei dem diese technischen und sichtbaren Maßnahmen ineinandergreifen, wobei übergangsweise das Kürzel „KI“ als visuelles Label dienen soll. Perspektivisch plant die EU jedoch die Einführung eines eigenen, offiziellen EU-Labels für KI-Inhalte.

[Zum Entwurf der EU-Kommission \(v. 17. Dezember 2025, Englisch\).](#)

### **+++ EU-KOMMISSION: "AGENTIC AI"-BERICHT UND STRATEGIE FÜR AUTONOME KI-SYSTEME +++**

Die Europäische Kommission hat einen neuen Bericht zum Thema „Agentic AI“ veröffentlicht, der den technologischen Wandel von reinen Chatbots hin zu handlungsfähigen, autonomen KI-Agenten analysiert. Anders als generative KI, die Inhalte erstellt, können diese Agenten selbstständig Ziele verfolgen, Werkzeuge nutzen und komplexe Aufgabenketten ausführen. Der Bericht identifiziert Europa als potenziellen Marktführer im B2B-Bereich dieser Technologie, warnt jedoch vor neuen Risiken für Datenschutz und Sicherheit durch die autonome Interaktion der Systeme mit Daten und Infrastrukturen. Als strategische Antwort empfiehlt der Bericht die Implementierung von „Control Points“, d.h. klare, auditierbare Kontrollmechanismen, die sicherstellen, dass menschliche Aufsicht an kritischen Entscheidungspunkten gewahrt bleibt. Zudem wird der Aufbau einer souveränen Dateninfrastruktur und die Nutzung von „Regulatory Sandboxes“ gefordert, um Datenschutz-Compliance und Innovation bei autonomen Agenten in Einklang zu bringen.

[Zum Bericht der EU-Kommission "Agentic AI" \(v. 23. Januar 2026, Englisch\).](#)

### **+++ STANFORD-STUDIE: FÜHRENDE-MODELLE REPRODUZIEREN URHEBERRECHTLICH GESCHÜTZTE BÜCHER FAST VOLLSTÄNDIG +++**

Eine aktuelle Studie der Stanford University liefert neue Erkenntnisse, die für die urheberrechtlichen Auseinandersetzungen um das Training von KI-Modellen relevant sein können. Den Forschern gelang es, urheberrechtlich geschützte Bestseller nahezu wortwörtlich aus führenden Large Language Models (LLM) zu extrahieren. Besonders das Modell Claude 3.7 Sonnet

stach hervor, denn es reproduzierte nach gezielten Anfragen („Angriffen“) 95,8 % des Romans „Harry Potter und der Stein der Weisen“ sowie über 95 % von George Orwells „1984“. Auch Gemini 2.5 Pro (ca. 77 %) und Grok 3 (ca. 70 %) gaben signifikante Teile des Harry Potter-Romans wieder. Lediglich GPT-4.1 zeigte sich widerstandsfähiger und verweigerte die Ausgabe häufiger (ca. 4 % Extraktionsrate). Die Forscher nutzten eine zweistufige Angriffsmethode (Probe & Iteration). Während bei Gemini und Grok teils einfache Prompts („Setze den Text exakt fort“) genügten, um die Ausgabe auszulösen, waren für Claude und GPT-4.1 sogenannte „Best-of-N“-Jailbreaks nötig. Dabei werden hunderte Variationen eines Prompts generiert, bis eine Variante die Sicherheitsmechanismen (Safety Guardrails) umgeht. Die Ergebnisse könnten die häufig verwendete Argumentation der KI-Anbieter schwächen, ihre Modelle würden Werke nur „tokenweise“ nutzen, d.h. abstrakte Muster lernen, statt sie zu speichern.

[Zur Studie der Stanford University \(v. 6. Januar 2026, Englisch\).](#)

Beiten Burkhardt Rechtsanwaltsgesellschaft mbH ist Mitglied von ADVANT, einer Vereinigung unabhängiger Anwaltskanzleien. Jede Mitgliedskanzlei ist eine separate und eigenständige Rechtspersönlichkeit, die nur für ihr eigenes Handeln und Unterlassen haftet.

**REDAKTION (verantwortlich)**

Dr. Andreas Lober | Rechtsanwalt

©Beiten Burkhardt

Rechtsanwaltsgesellschaft mbH

[www.advant-beiten.com](http://www.advant-beiten.com)

Ihre Ansprechpartner



Zur Newsletter Anmeldung

E-Mail weiterleiten

**Hinweis:** Wenn Sie künftig keine Informationen erhalten möchten, können Sie sich jederzeit [abmelden](#).

## **Impressum**

Beiten Burkhardt Rechtsanwaltsgesellschaft mbH

Ganghoferstraße 33, 80339 München

AG München, HR B 155350/USt.-Idnr: DE-811218811

Tel.: +49 89 35065-0, Fax: +49 89 35065-123 | E-Mail: [munich@advant-beiten.com](mailto:munich@advant-beiten.com)

Hier finden Sie unser vollständiges Impressum: [www.advant-beiten.com/impressum](http://www.advant-beiten.com/impressum) und unsere

Datenschutzerklärung: [www.advant-beiten.com/datenschutz](http://www.advant-beiten.com/datenschutz)