



News, Tipps und Debatten zu rechtlichen Fragen rund um generative KI

Oktober 2025

Liebe Leser*innen,

der Herbst ist da, die Tage werden kürzer – und unser Newsletter dieses Mal auch. Nichtsdestotrotz haben wir einige spannende Beiträge für Sie im Angebot.

Im Interview mit der Wirtschaftsinformatikerin und Rechtsanwältin Lisa Käde sprechen wir über **KI-Plagiate, die Pastiche-Regelung und – ja – Besonderheiten der Popmusik**.

Außerdem gibt es mal wieder **Neues zur KI-Verordnung**: Die EU-Kommission lädt dazu ein, sich in die Ausgestaltung einiger Praxisleitfäden (Codes of Conduct) einzubringen. Diesmal geht es um die Themen **Kennzeichnung, Information und Transparenz**. Wir erklären, wer mitmachen kann und was es zu bewirken gibt.

Transparenz wird auch dann zum Thema, wenn Angestellte im Rahmen ihrer Arbeit KI-Anwendungen einsetzen, die nicht offiziell freigegeben sind. Dann spricht man auch von **Schatten-KI**. Was genau sich dahinter verbirgt und welche rechtlichen Probleme damit verbunden sind, haben wir uns für Sie angesehen.

Wir wünschen Ihnen, wie immer, eine erhellende Lektüre,
Ihr **prompt/-Team**

Wenn Ihnen unser Newsletter gefällt, sagen Sie's gerne weiter – Dankeschön :-)

prompt/ empfehlen

Wenn Sie **prompt/** regelmäßig bekommen möchten, können Sie sich über diesen Button anmelden:

prompt/ bestellen

Ebenso freuen wir uns über **Themenvorschläge, Fragen oder Erfahrungen aus Ihrem beruflichen oder persönlichen Alltag mit generativer KI**.

Schreiben Sie uns eine E-Mail an prompt@irights-lab.de

/Inhalt

- [Interview](#) | Lisa Käde über KI-Outputs und die Suche nach einer persönlichen geistigen Schöpfung
- [KI-Verordnung](#) | Konsultation der EU-Kommission: Mitwirkung bei Praxisleitfäden zu Kennzeichnung, Information und Transparenz
- [Hintergrund](#) | Was ist Schatten-KI?
- [Weiterlesen](#) | Cloudflare vs. KI-Bots, neue Haftungsregeln für Software und KI, Perplexity führt Bezahlmodell für Verlage ein ... und mehr!
- [Über uns](#) | [Impressum](#)

/Nachfrage

KI-Outputs und die Suche nach einer persönlichen geistigen Schöpfung

Kopie, Plagiat oder künstlerische Bearbeitung? Wenn generative KI Werke ausgibt, die vorhandenen Originalen ähneln, stellen sich urheberrechtliche Fragen. Lisa Käde über das Konzept der „unabhängigen Doppelschöpfung“, über die „Pastiche“-Regelung und die Frage nach den Intentionen.



[Lisa Käde](#) ist Wirtschaftsinformatikerin und Rechtsanwältin in der Kanzlei JBB Rechtsanwälte:innen in Berlin und berät in allen rechtlichen Fragen rund um Entwicklung und Nutzung von KI-Systemen sowie Open Source-Compliance.

Foto: JBB Rechtsanwälte Partnerschaft

prompt/: In Bezug auf den Output von generativer KI zitierten Sie, „wenn Mickey

rauskommt, muss Mickey auch drin sein ... “ – was ist damit gemeint?

Lisa Käde: Dahinter steckt die verbreitete Annahme, dass ein Prompt in irgend einer Weise zwangsläufig zur Reproduktion führt – ähnlich dem Abrufen von Daten aus einer Datenbank. Dabei wird aber verkannt, dass generative KI-Systeme nicht einfach Gelerntes wiedergeben können oder sollen. Diese Systeme lernen vor allem, zu verallgemeinern und mithilfe von Statistik die „Essenz“ des Gelernten wiederzugeben. Außerdem kann ein User, der es darauf anlegt, möglicherweise durch eine extrem präzise und ausführliche Formulierung des Prompts das KI-System zu einem Output leiten, der sehr nah an bestimmten Trainingsdaten liegt – weil es etwa in Bezug auf bestimmte Merkmale nichts zu verallgemeinern gibt. Dementsprechend handelt es sich bei dem generierten Output nicht um eine „Kopie“ im technischen Sinne, denn es wird nichts einfach „wiedergegeben“, sondern es wird ein neuer Output errechnet.

prompt/: Ist es urheberrechtlich nicht unerheblich, wie ein Plagiat zustande kommt, es ist so oder so eine Urheberrechtsverletzung – oder?

Lisa Käde: Jein. Es gibt das Konzept der „unabhängigen [Doppelschöpfung](#)“: Wenn zwei Menschen die gleiche Idee haben und sie auf gleiche Weise umsetzen und ein sehr ähnliches oder gleiches Werk dabei herauskommt, ohne dass beide voneinander abgekupfert haben können, kann für beide Urheberrechtsschutz entstehen. Es wurde auch schon für generative KI-Modelle angedacht, ob ein derartiger Ansatz bei der urheberrechtlichen Betrachtung des Outputs zu berücksichtigen ist – insbesondere wenn das, was scheinbar plagiiert wurde, nicht in den Trainingsdaten enthalten war. Womit ich Ihnen aber Recht gebe: Es kommt nicht darauf an, wie der Output – technisch – zustande gekommen ist. Wenn der Output Trainingswerken stark ähnelt oder gleicht, kann eine Vervielfältigung vorliegen – egal ob aus einer Datenbank abgerufen oder statistisch aufgrund des Prompts neu erzeugt.

prompt/: Lassen sich diese Überlegungen auf die Klage der GEMA gegen das Unternehmen Suno anwenden? Die GEMA legte Beispiele dafür vor, dass die Suno-KI Werke erzeugt, die sehr nah an Originalen lägen und damit als unerlaubte Vervielfältigungen zu sehen seien.

Lisa Käde: In der Pop-Musik kommt vermutlich verschärfend dazu, dass sich einige Muster etabliert haben, die sich in vielen Werken wiederfinden. Beispielsweise bestimmte Akkordfolgen – Am, F, C, G – bestimmte Rhythmen oder Song-Strukturen, sodass allein in den Trainingsdaten – also in dem Musikbestand – schon vieles ähnlich klingen dürfte. Da auch Musik-KI darauf getrimmt ist, zu produzieren, was gefällt, liegt es nahe, dass die Ergebnisse in diese bekannten und bewährten Muster fallen und dementsprechend ähnlich klingen. Zudem vermute ich, dass der Bestand aktueller Pop-Musik – aus einer bestimmten zeitlichen Epoche – deutlich eingeschränkter sein dürfte als der Bildbestand und dass auch die Parameter deutlich eingeschränkter sind (Takt, Tempo, Tonart, Rhythmus, Dynamik, Instrument, Stimme, Text). Auch deshalb können hier größere Ähnlichkeiten bestehen.

prompt/: Im EU-Urheberrecht gibt es die „Pastiche“-Regelung, wonach künstlerisch verfremdende Bearbeitungen geschützter Werke erlaubt sein können. Wäre das ein rechtlicher Weg für solche KI-generierten Werke, in denen Originale erkennbar, aber eben nicht vervielfältigt sind?

Lisa Käde: Hier dürfte es unter anderem darauf ankommen, ob die Nutzerin es mit ihrem Prompt darauf anlegt, verschiedene Werke zusammenzuführen oder fremde Werke zu nutzen. So könnte sie sich etwa Kombinationen überlegen wie „Mona Lisa spaziert durch Paris“ oder „Mickey und Donald stehen in Monet's Garten“. Wenn eine Übernahme fremder Werkteile bewusst und im Rahmen einer geistigen Auseinandersetzung erfolgt, könnte die Pastiche-Schranke auch in diesem Fall an Bedeutung gewinnen – wobei auch noch abzuwarten ist, wie sich der Pastiche-Begriff in der Rechtsprechung entwickelt. Spannender ist vielleicht sogar, wie es aussieht, wenn eine Übernahme gar nicht beabsichtigt war. Etwa durch den Prompt „zwei Enten stehen in einem malerischen Garten“, aber dieser Prompt dennoch durch das KI-System so interpretiert wird, dass im Ergebnis eine Übernahme fremder Werkteile erfolgt. Dann dürfte es überaus fraglich sein, ob eine inhaltliche beziehungsweise geistige Auseinandersetzung mit dem Ausgangswerk erfolgt, die die Anwendung der Pastiche-Schranke rechtfertigt. Auch beobachten wir, dass es nicht ausreicht – wie früher, „vor KI“ – nur auf das Ergebnis zu schauen, sondern verstärkt auch auf den dahinter liegenden Prozess, die „Werkschöpfungskette“. Im Rahmen der Ermittlung des Urheberrechtsschutzes an KI-Output passiert das jetzt schon auf der Suche nach einer persönlichen geistigen Schöpfung: Es wird untersucht, ob ein Mensch das KI-System so gesteuert hat, dass der Output eine persönliche geistige Schöpfung darstellt. Und auch beim Pastiche wird vermutlich zu erforschen sein, ob es „die Idee der KI war“, bestimmte Elemente zu nutzen und zu kombinieren, oder ob ein Mensch das beabsichtigt hat.

/KI-Verordnung

Konsultation der EU-Kommission: Mitwirkung bei Praxisleitfäden zu Kennzeichnung, Information und Transparenz

Anfang September informierte die Europäische Kommission über eine weitere öffentliche Konsultation. Sie bezieht sich auf weitere Praxisleitfäden (englisch: Codes of Practice) zur Umsetzung der KI-Verordnung (KI-VO). Bis zum 02. Oktober konnten Interessierte sich [bei der Kommission melden](#), um am „Code of Practice on Transparent AI Systems“ mitzuwirken. Im Mittelpunkt der neuen Arbeiten stehen die sogenannten Transparenzvorschriften des Artikel 50 der KI-VO. Die Ergebnisse der Konsultation sollen die EU-Kommission darin unterstützen, die enthalten Kennzeichnungsregeln zu präzisieren

KI-VO-Kennzeichnungspflichten für Deep Fakes, KI-Outputs, bestimmte Chatbots und KI-Agenten

Im genannten Artikel 50 sind unterschiedliche Informations- und Kennzeichnungspflichten zusammengefasst. Manche betreffen gewerbliche Nutzer*innen. Sie werden beispielsweise dazu verpflichtet, selbst produzierte Deepfakes oder automatisiert produzierte Nachrichtentexte als solche zu kennzeichnen. Unternehmen, die generative KI-Systeme entwickeln, müssen unter anderem dafür sorgen, dass der von ihren Systemen produzierte Output als künstlich erzeugt oder als KI-manipuliert erkennbar ist. Zudem soll diese Kennzeichnung als sogenanntes „digitales Wasserzeichen“ maschinenlesbar sein. KI-Systeme, die direkt mit Menschen interagieren – etwa Chatbots oder entsprechende KI-Agenten (siehe hierzu [prompt/ Juli 2025](#)) – müssen als solche erkennbar ausgewiesen werden.

Kennzeichnungsregeln sind noch zu unpräzise

Wie so häufig bei Verordnungen sind auch im Artikel 50 der KI-VO einige Details zur Umsetzung der Pflichten noch unklar: Wie muss zum Beispiel eine Kennzeichnung gestaltet sein, damit ein Chatbot klar und eindeutig als solcher erkennbar ist? Welche technischen Merkmale muss ein digitales Wasserzeichen aufweisen, damit es interoperabel, belastbar und zuverlässig wirksam ist, wie es der Artikel 50 in Absatz 2 fordert? Auf diese und weitere Fragen sollen die angekündigten Praxisleitfäden eingehen und den Hersteller*innen, Anbieter*innen, Betreiber*innen und Verbraucher*innen praxisnahe Hilfestellungen bieten.

Praxisleitfäden sollen in der ersten Jahreshälfte 2026 vorliegen

Auch diesmal – wie schon bei den bereits vorliegenden Praxisleitfäden für generative KI-Modelle (siehe [prompt/ August 2025](#)) – dürfen Unternehmen, Verbände, wissenschaftliche Einrichtungen und weitere Stakeholder ihre Bedarfe und Vorschläge im Rahmen der Konsultation einbringen. Alle Eingaben werden zusammen mit Vertreter*innen aus Wirtschaft, Wissenschaft und Zivilgesellschaft in Arbeitsgruppen ausgewertet, die daraufhin Entwürfe erarbeiten. Die fertigen Praxisleitfäden sollen laut EU-Kommission in der ersten Jahreshälfte 2026 vorliegen – rechtzeitig, bevor die erwähnten Kennzeichnungspflichten der KI-Verordnung ab August 2026 in Kraft treten. Die in der KI-VO verankerten Regeln sind verbindlich einzuhalten. Demgegenüber positioniert die EU die Praxisleitfäden als rechtlich nicht verbindliche Umsetzung- und Interpretationshilfen (auch dazu mehr in [prompt/ August 2025](#)).

Was ist Schatten-KI?

Was mystisch klingt, ist eigentlich ganz einfach: Wenn am Arbeitsplatz nicht genehmigte KI-Anwendungen genutzt werden, spricht man von Schatten-KI. Ein Beispiel: Eine Schule besitzt keine offizielle Lizenz, die Lehrenden nutzen aber – ohne Erlaubnis – ihren privaten ChatGPT-Account, um Aufgaben zu erstellen oder Arbeiten ihrer Schüler*innen zu korrigieren. Oder: Eine Mitarbeiterin in der Verwaltung muss einen Sachbericht erstellen und gibt dafür Daten in einen privat abonnierten Chatbot ein. Jetzt könnte man fragen: Wo ist das Problem?

Generell geht es Arbeitgebern zunächst darum, dass die von den Beschäftigten genutzten Werkzeuge und Materialien legal erworben, funktional geeignet und geprüft sind, auch um damit interne Vorgaben und Qualitätsmaßstäbe einhalten zu können. Bei KI-Tools meint das unter anderem den Umgang mit sensiblen Firmendaten oder personenbezogenen Daten. Verwenden die Beschäftigten nun eigenmächtig Anwendungen – im Schatten der Kontrollmechanismen – kann dies unerwünschte Folgen haben: Im Fall KI-gestützter Chatbots in der Cloud besteht die Gefahr, dass sensible Daten unkontrolliert an die KI-Anbieter abfließen. Oder dass KI-generierte Outputs fehlerhaft sind und womöglich falsche beziehungsweise erfundene Informationen enthalten. Bei daraus resultierenden Schadensfällen könnte der Arbeitgeber dafür haftbar gemacht werden.

Was kann man dagegen tun?

Die einfache Antwort wäre vermutlich: Lizenzen anschaffen, Konfigurationen steuern und die Nutzung möglichst allen zugänglich machen – auch externen, vertraglich eingebundenen Mitarbeiter*innen. Wenn das nicht möglich ist, ließe sich vielleicht per anonymer Umfrage ermitteln, welche KI-Tools von den Beschäftigten genutzt beziehungsweise gefordert werden. Daraufhin können die Verantwortlichen Risikoeinschätzungen vornehmen und Richtlinien für den Einsatz der ausgewählten KI-Anwendungen erstellen. Darin kann auch definiert sein, was ausdrücklich *nicht* erlaubt ist. Nicht zuletzt sollten Mitarbeitende im Zuge einer organisierten KI-Einführung qualifiziert werden, wie es die [KI-Verordnung ohnehin verpflichtend vorsieht](#).

/Weiterlesen

Artikel und Quellen zu aktuellen Rechtsfragen bei generativer KI

- **Cloudflare will KI-Bots blockieren:** Wie das Unternehmen Verlagen und

Webseitenbetreibern helfen möchte, ihre Inhalte vor dem ungewollten Abgreifen durch KI-Crawler zu schützen. <https://taz.de/Wie-Cloudflare-KI-Bots-aufhalten-will/!6111259/>

- **„Rolling Stone“-Verlag verklagt Google wegen KI-Zusammenfassungen:** Der US-Medienkonzern Penske Media zieht vor Gericht, weil Google journalistische Inhalte ohne Zustimmung für KI-generierte Suchergebnisse nutzt – und damit Traffic und Einnahmen von Verlagen abzieht. <https://www.zeit.de/digital/2025-09/rolling-stone-klage-ki-zusammenfassungen-google>
- **Neue Haftungsregeln für Software und KI:** Das deutsche Justizministerium plant, Hersteller von Software und KI für Schäden haften zu lassen – etwa bei Unfällen mit autonomen Fahrzeugen oder fehlerhaften Updates. <https://www.heise.de/news/Justizministerium-So-sollen-Hersteller-von-Software-und-KI-fuer-Produkte-haften-10641780.html>
- **Bundesjustizministerium: Entwurf zur Produkthaftung für KI und Software** (offizielle Seite) https://www.bmjv.de/SharedDocs/Gesetzgebungsverfahren/DE/2025_Produkthaftung.html
- **Perplexity AI startet Bezahlmodell für Verlage:** Die KI-Suchmaschine führt ein Abo-Modell ein, das Verlage direkt an den Einnahmen beteiligt – ein neuer Ansatz, um Journalismus im KI-Zeitalter zu finanzieren. <https://www.bdzv.de/service/presse/branchennachrichten/2025/perplexity-ai-fuehrt-bezahlmodell-fuer-verlage-ein>

/Über uns | Impressum

Diesen Newsletter erstellen und produzieren wir, das [iRights.Lab](#), im Rahmen des Forschungsprojekts [Generative KI – Innovation und Recht in Arbeitsprozessen \(GenKI-IR\)](#), gefördert vom Bundesministerium für Forschung, Technologie und Raumfahrt (BMFTR).

Das Projekt hat zum Ziel, rechtliche und gesellschaftliche Rahmen zu beschreiben, mit denen Anwendungen generativer KI unterschiedliche Arbeitsprozesse sinnvoll unterstützen können.

Gefördert durch:



Bundesministerium
für Forschung, Technologie
und Raumfahrt

Texte: Merlin Münch, Dr. Matthieu Binder, Henry Steinhau, Ella Jordan

Redaktion: Henry Steinhau (verantwortlich), Merlin Münch

Mitarbeit und Lektorat: Solvejg Gunkel, Ella Jordan

Rechtehinweis: Alle Texte und Bilder stehen unter der offenen Lizenz [CC-BY 4.0](#),
außer das Porträts von Lisa Käde, Foto: JBB Rechtsanwälte Partnerschaft mbB.

iRights.Lab GmbH

www.iriights-lab.de

Oranienstr. 185

10999 Berlin

prompt@iriights-lab.de

Tel: +49 30 40 36 77 230

Fax: +49 30 40 36 77 260



Steuernr. 37/359/5262 | UStID DE311181302

Registergericht:

Amtsgericht Charlottenburg Nr. HRB 185640 B

Geschäftsführer*innen:

Philipp Otto, Dr. Wiebke Glässer

Diese E-Mail wurde an hest@hest.de gesendet.
Sie haben diese E-Mail erhalten, weil Sie den prompt/-Newsletter bestellt haben.

[Abbestellen](#)