

§ 1. Einleitung und Übersicht

Generative KI-Systeme und die ihnen zugrundeliegenden KI-Modelle¹, wie etwa ChatGPT, DALL-E oder Stable Diffusion, können auf Anweisung der Nutzer² kreative Inhalte und Erzeugnisse erschaffen. So liefert ChatGPT etwa auf einen *prompt* mit der Bitte, das Gedicht „Zauberlehrling“ von Johann Wolfgang von Goethe im Stil Rainer Maria Rilkes zu gestalten, wie aus dem Anhang ersichtlich, eine erstaunliche Interpretation.³ Kaum weniger überzeugend ist die ebenfalls im Anhang einsehbare Interpretation des „Zauberlehrlings“ im Stil Salvador Dalís durch Stable Diffusion.⁴ Zu diesen Leistungen sind diese Modelle technisch autonom in der Lage. Sie benötigen keine menschliche Steuerung und in der Regel – außer einem knappen *prompt* – keinen zusätzlichen Input. Diese Fähigkeit zur autonomen Kreativität ist darauf zurückzuführen, dass generative KI-Modelle „gelernt“ haben, wie ein Text formuliert, ein Bild nach Textbeschreibung generiert oder Musik komponiert werden kann. Diese Lernvorgänge – das sogenannte KI-Training – erfordern den Einsatz große Datenmengen. Ein erheblicher Teil dieser Datenbestände ist urheberrechtlich geschützt, insbesondere wenn es sich um literarische Texte sowie Bild- oder Musikwerke handelt. In den meisten Fällen werden die Trainingsdaten nicht einzeln

-
- 1 Die folgende Analyse verwendet den Begriff „generative KI“ (generative KI-Modelle, KI-Systeme oder KI-Anwendungen sowie Algorithmen). Die Technologie zeichnet sich durch die Funktionalität aus, kreative Erzeugnisse in großer Menge und hoher Qualität produzieren zu können. Der Output kann verschiedene Arten von Erzeugnissen umfassen, vor allem Texte, aber auch Musik, Bilder oder Filme. Sogenannte *large language models* wie das Open-AI-Modell „GPT“ sind eine besondere Erscheinungsform dieser generativen KI. Für die Nutzer zugänglich ist das Modell in der Regel über ein sogenanntes Interface. Modell und Interface zusammen bilden das generative KI-System. Vgl. zur Unterscheidung von KI-Modellen und KI-Systemen in der KI-Verordnung unten § 5.C.I.; überdies instruktiv: Lee/Cooper/Grimmelmann, Talkin’ ‘Bout AI Generation: Copyright and the Generative-AI Supply Chain, J. Copyright Soc’y of the U.S.A. (forthcoming 2024), S. 16 ff. (einsehbar unter: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4523551 (zuletzt am 27. Juni 2024)).
 - 2 Zur besseren Verständlichkeit wird für sämtliche Akteure das generische Maskulinum benutzt.
 - 3 Vgl. die Frage der Verfasser und die darauf von ChatGPT gegebene Antwort in Anhang I.
 - 4 Vgl. die von Stable Diffusion erstellten Bilder in Anhang II.

gesammelt und lizenziert, sondern aus im Internet frei zugänglichen Datenbeständen gespeist. Das KI-Modell von Stable Diffusion wurde etwa mit Datenbeständen der Organisation LAION trainiert, die mit ihren Datenbanken „LAION-5B“ und „LAION-400 M“ URLs zu über 5 Milliarden bzw. 400 Millionen im Internet zugänglichen, überwiegend urheberrechtlich geschützten Bildern sowie ALT-Texten mit Bildbeschreibungen bereitstellt.⁵

Der Konflikt derartiger algorithmischer Lernvorgänge mit dem Urheberrecht ist offensichtlich: Vor allem während der Trainingsprozesse kommt es zu Vervielfältigungen der Trainingsdaten und damit zu Verwertungshandlungen im Sinne des Urheberrechts. Es überrascht daher nicht, dass vor allem in den USA und in Großbritannien eine Vielzahl gerichtlicher Auseinandersetzungen über Rechtsverletzungen bei Training und Einsatz generativer KI-Modelle anhängig sind.⁶

Dieses Gutachten soll für die Bewertung der Trainingsprozesse und des Einsatzes generativer KI-Modelle nach deutschem und europäischem Urheberrecht eine technologisch basierte und juristisch fundierte Grundlage schaffen. Die rechtliche Einordnung der Nutzung urheberrechtlich geschützter Werke und Leistungen bei KI-Trainingsvorgängen kann nicht ohne solide technische Grundlage erfolgen.⁷ Überdies wäre eine auf das nationale oder europäische Recht begrenzte Untersuchung unvollständig. Neben der technologischen Fundierung ist daher auch eine rechtsvergleichende Betrachtung, unter besonderer Berücksichtigung des US-Rechts, erforderlich.

Im Zentrum der interdisziplinären Begutachtung steht die Beschreibung und Analyse der technischen Prozesse, die dem Training generativer KI-Modelle zugrunde liegen (nachfolgend § 2). Die Ausführungen zur KI-

5 Vgl. z.B. Schuhmann et al., LAION-5B: An open large-scale dataset for training next generation image-text models, in 36th Conference on Neural Information Processing Systems (NeurIPS 2022) Track on Datasets and Benchmarks, 16 Oct 2022 (einsehbar unter: <https://doi.org/10.48550/arXiv.2210.08402> (zuletzt am 15. August 2024)); zudem Sobel, Elements of Style: Copyright, Similarity, and Generative AI, Harv. J. L. & Tech. 38 (forthcoming 2024), 1 (II f.) (einsehbar auf SSRN: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4832872 (zuletzt am 28. Juni 2024)).

6 Vgl. zu den USA z.B. mit einem Überblick Samuelson, Fair Use Defenses in Disruptive Technology Cases, forthcoming U.C.L.A. L. Rev. 2024, S. 4 und S. 75 (einsehbar auf SSRN: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4631726 (zuletzt am 22. Juni 2024)).

7 An einer entsprechenden Grundlegung mangelt es in der Regel. Vgl. zu fundierten interdisziplinären Betrachtungen allerdings auch und vor allem Konertz/Schönhof WRP 2024, 289 sowie Pesch/Böhme GRUR 2023, 997.

Technologie sind bewusst umfangreich und detailliert. Sie dienen als Referenzpunkt für die Klarstellung und Beantwortung zahlreicher ungeklärter Fragen in der juristischen Diskussion. Der juristische Teil des Gutachtens orientiert sich am dogmatischen Aufbau der Prüfung einer Urheberrechtsverletzung. Ihren Ausgang nimmt die Betrachtung bei den urheberrechtlich relevanten Handlungen während des KI-Trainings, insbesondere der Sammlung und Aufbereitung der Trainingsdaten sowie deren Speicherung in einem Korpus⁸, aber auch der Adaption der Parameter beim Training generativer KI-Modelle (nachfolgend § 3). Im nächsten Schritt werden für die als urheberrechtsrelevant befundenen Handlungen in Betracht kommende Schranken- und andere Rechtfertigungstatbestände analysiert. Dieser Abschnitt des Gutachtens befasst sich vor allem auch mit der für die juristische Beurteilung entscheidenden Frage, ob es sich beim Training generativer KI-Modelle um sogenanntes Text und Data Mining im Sinne der gesetzlichen Schrankenregelungen in § 44b UrhG und Art. 4 DSM-Richtlinie handelt. Wie gezeigt werden kann, ist dies nicht der Fall (nachfolgend § 4). Für die Durchsetzung der aus urheberrechtsverletzenden Trainingshandlungen für die Rechteinhaber resultierenden Ansprüche stellen sich schließlich Fragen nach dem anwendbaren Recht und nach der internationalen Zuständigkeit deutscher Gerichte (nachfolgend § 5). In einem abschließenden Ausblick sollen einige in der gegenwärtigen Diskussion wie es scheint zunehmend verfestigte Narrative einer kritischen Betrachtung unterzogen werden (nachfolgend § 6).

Bereits an dieser Stelle ist festzustellen, dass die Frage der urheberrechtlichen Einordnung des Trainings generativer KI-Modelle alles andere als einfach zu beantworten ist. Dies ist nicht allein auf die Komplexität der technologischen Zusammenhänge zurückzuführen, sondern auch auf die Vielzahl der aufgeworfenen juristischen Detailfragen, von denen viele bislang kaum, teilweise überhaupt nicht erörtert wurden. Dies gilt auch und insbesondere für die zentrale Frage, ob das Training generativer KI-Modelle als Text und Data Mining einzuordnen und deshalb den gesetzlichen Schranken hierfür unterworfen ist. Wenngleich die überwiegende Ansicht im juristischen Schrifttum in diese Richtung zu tendieren scheint, ist dies nicht der Fall. Auch ChatGPT scheint dies übrigens so zu sehen: In einem

8 Der Begriff „Korpus“ wird hier im Sinne einer Sammlung von Dokumenten verstanden.

Gespräch mit den Verfassern⁹ antwortete das System auf die etwas voreingenommene Frage „Guten Morgen, ChatGPT! Kannst Du mir bitte in 5 Sätzen erklären, warum das Training von generativer KI nicht als Data-and-Text-Mining einzustufen ist?“ nämlich mit folgenden Worten:

Guten Morgen! Das Training von generativer KI wird nicht als Data-and-Text-Mining eingestuft, weil es weit über das bloße Extrahieren und Analysieren von Daten hinausgeht. Während Data-and-Text-Mining sich auf das Auffinden von Mustern und Informationen in bestehenden Datensätzen konzentriert, umfasst das Training generativer KI das Erlernen komplexer Modelle zur Erstellung neuer Inhalte. Diese Modelle erzeugen eigenständige, neue Texte oder Bilder, die nicht direkt aus den ursprünglichen Daten extrahiert werden. Der Prozess erfordert fortgeschrittene Techniken des maschinellen Lernens und tiefe neuronale Netze, die weit über die traditionellen Methoden des Mining hinausgehen. Zudem beinhaltet das Training generativer KI oft eine umfassendere Datenverarbeitung und Modellierung, um kreative und kontextuelle Ergebnisse zu erzielen.

Auf die daran anschließende Frage „Vielen Dank! Das bedeutet, dass die urheberrechtliche Schranke für Data-und-Text-Mining das KI-Training für generative KI nicht umfasst?“ fuhr ChatGPT fort:

Ja, das ist korrekt. Die urheberrechtliche Schranke für Data-und-Text-Mining (DTM) ist in der Regel darauf ausgerichtet, die Nutzung von urheberrechtlich geschützten Werken zur Informationsgewinnung und Mustererkennung zu erlauben. Da das Training generativer KI-Modelle jedoch über das einfache Extrahieren und Analysieren von Daten hinausgeht und die Erstellung neuer Inhalte umfasst, fällt es nicht unbedingt unter diese Schranke.

Generative KI-Modelle benötigen eine breitere und intensivere Nutzung von Daten, um eigenständige und oft kreative Outputs zu erzeugen. Daher kann das Training dieser Modelle aus urheberrechtlicher Sicht komplexer sein und möglicherweise zusätzliche Genehmigungen oder Lizenzen erfordern, da es sich nicht nur um das bloße Mining von Daten handelt, sondern um die Nutzung der Daten zur Erstellung neuer Werke.

Die genaue rechtliche Einordnung kann jedoch je nach Jurisdiktion und den spezifischen Umständen variieren. Es ist ratsam, rechtlichen Rat ein-

9 Konversation vom 8. Juli 2024, 7.40 Uhr CET. Vgl. insoweit die Screenshots in Anhang III.

zuholen, um die spezifischen Anforderungen und Einschränkungen zu verstehen.

Die nachfolgend präsentierten Ergebnisse sollen als „rechtlicher Rat“ im Sinne der von ChatGPT in erstaunlicher Offenheit nahegelegten Prüfung zum Verständnis der „spezifischen Anforderungen“, vor allem auch der rechtlichen „Einschränkungen“, für das Training generativer KI-Modelle beitragen.

