

Lokale KI-Anwendung im Bewegtbild-Studio

Transparente, souveräne und replizierbare KI-Praxis an Kunsthochschulen lehren. Ohne Cloud-Abhängigkeit.

Antrag von:

Ting-Chun Liu

Künstlerischer Mitarbeiter (Creative Technologist)
Professur Crossmediales Bewegtbild

Visuelle Kommunikation, Fakultät Kunst und Gestaltung
Bauhaus-Universität Weimar

Zusammenfassung

Da der Einfluss Künstlicher Intelligenz im kreativen Feld zu einer wiederkehrenden Frage unter Studierenden geworden ist, unternimmt dieser Antrag den Versuch, ein didaktisches Experiment für die Hochschullehre der Fakultät Kunst und Gestaltung zu entwickeln, das eine Zusammenarbeit mit dieser Technologie ermöglicht und dabei eine informierte und kritische Haltung wahrt. Studierende im Bewegtbild arbeiten täglich mit Bildern. Doch was sieht eine Maschine, wenn sie dasselbe Material betrachtet? Wie beschreibt eine Maschine eine Szene mit emotionalem Gewicht? Welche Gesten übersieht sie systematisch? Können Studierende ihr Filmmaterial nach Atmosphäre oder Emotionen analysieren lassen? Dieses Lehrprojekt rückt diese Frage in den Mittelpunkt der Professur für crossmediales Bewegtbild an der Bauhaus-Universität Weimar.

Das Projekt führt KI-gestützte Bildwahrnehmung als eine konkrete, lokal betriebene Bildpraxis ein. Es werden drei Open-Source-Tools eingesetzt, um zu untersuchen, welche alternativen Erzählungen entstehen, wenn eine Maschine Filmmaterial analysiert: Ollama als lokale Laufzeitumgebung, LLaVA zur Übersetzung von Einzelbildern in Sprache und CLIP zum Abgleich von Bildern und Text. Alle Modelle laufen auf der Hardware der Studierenden. Keine Cloud, keine API-Schlüssel, keine Datenübertragung. Der Prozess geschieht im lokalen Raum, sichtbar, messbar, anfechtbar. Der lokale Betrieb der Modelle ohne Cloud-Abhängigkeit macht zudem die materiellen Bedingungen Künstlicher Intelligenz greifbar. Eine Feldarbeit im Rechenzentrum erweitert diesen Blick vom Laptop-Lüfter bis zum Kühlturm und verortet die alltägliche Nutzung von Computation in der Kreativwirtschaft in ihrem industriellen und ökologischen Maßstab.

1. Persönliche Motivation

Ich arbeite als Künstlerischer Mitarbeiter an der Professur Crossmediales Bewegtbild der Bauhaus-Universität Weimar und lehre Creative Coding, interaktive Medien sowie kritische Künstliche Intelligenz in der künstlerischen Praxis. Parallel dazu pflege ich eine aktive Forschungspraxis, die den materiellen und politischen Bedingungen nachgeht, welche die Ästhetik technologischer Systeme prägen. Wie Simon Penny formulierte: „Künstler, die sich mit diesen Technologien beschäftigen, setzen sich gleichzeitig auch mit der Ökonomie der Konsumgüter auseinander“ (Penny, 1995). Meine Forschung stellt einen anhaltenden Versuch dar, diese Verflechtung lesbar zu machen, sowohl in der Lehre als auch in der Praxis. Ich habe eine Arbeit zur Stereotypen-Kodierung in großen Vision-Language-Modellen vorgestellt und gemeinsam mit Leon-Etienne Kühr eine Videoarbeit entwickelt, die KI-Bild-Rückkopplungsschleifen erforscht (gezeigt im Museum Angewandte Kunst Frankfurt, 2026).

Im Seminaralltag begegne ich Studierenden, die KI-Werkzeuge selbstverständlich nutzen, ohne zu verstehen, wie ihre Ergebnisse zustande kommen. Im Umgang mit diesen Werkzeugen habe ich aus erster Hand erfahren, was ich in meinem Essay als Rückkopplungsschleife beschrieben habe: „as the software changes due to the need for better quality for further commercializing, the expectations also shift gradually. It limits how we perceive tools and what kind of output we are able to generate“. Die Studierenden befinden sich in dieser Schleife, ohne es zu wissen. Die kommerziellen KI-Schnittstellen, die Sie täglich nutzen, sind speziell darauf ausgelegt, ihre eigenen Funktionsmechanismen zu verbergen. Dieses Stipendium bietet die Möglichkeit, genau das Gegenteil zu tun. Ziel ist es, ein Unterrichtsformat zu entwickeln, das es den Studierenden durch den praktischen Einsatz von KI ermöglicht, aus erster Hand zu erfahren, wie KI-Modelle funktionieren und aufgebaut sind.

2. Problemstellung

2.1 Die veränderte Situation im Bewegtbild

Automatisierte Bilderkennung ist längst keine Zukunftstechnologie in der Bewegtbild-Produktion mehr. Streaming-Plattformen versehen Filmmaterial automatisch mit Schlagworten, Schnittsoftware schlägt Schnittpunkte vor, KI-Systeme erzeugen Untertitel aus Bildinhalten. Die Funktion des Bildes hat sich von der Darstellung hin zur Aktivierung und zu Operationen verlagert (Paglen, 2016). Dies wirft für Studierende des Bewegtbilds eine neue Frage auf: nicht nur, was ein Bild zeigt, sondern auch, was eine Maschine darin liest. Hinter diesen Funktionen stehen zwei Modellklassen, die im Zentrum dieses Seminars stehen. Die erste sind Vision-Language-Modelle, also Systeme, die visuellen Input in Textbeschreibung übersetzen. Ausgehend von einem Videoframe erzeugen sie einen Satz: was sie sehen, wer anwesend ist, was geschieht. Die zweite sind kontrastive Bild-Text-Modelle, die Bilder und Text in einem gemeinsamen semantischen Raum verorten. Sie bilden zugleich die Grundlage der meisten Bildgenerierungsprozesse und ermöglichen die Suche nach Beschreibung statt nach Schlagwort. Beide Modellklassen sind inzwischen ausgereift, quelloffen und auf lokaler Hardware ohne Cloud-Zugang lauffähig.

Kommerzielle Anbieter haben kein Interesse daran, diese Systeme lesbar zu machen. KI-Bilderkennungssysteme sind Blackboxes hinter Bezahlschranken. Ihre Ergebnisse lassen sich akzeptieren oder zurückweisen, aber nicht verstehen, befragen oder umlenken. Der vorliegende Antrag soll dazu beitragen den Studierenden einen kritischen Umgang mit den Systemen zu lehren. Dafür muss die Lehre tiefer ansetzen.

2.2 Warum lokale Modelle?

Die übliche KI-Didaktik greift zur Cloud: ein Online-Jupyter-Notebook, ein API-Zugangsschlüssel, ein gehostetes Modell. Das ist nachvollziehbar, denn Cloud-Dienste sind schnell und liefern hochwertige Ergebnisse. Allerdings übersieht diese Arbeitsweise die didaktisch wichtigsten Fragen. Wenn Bilderkennung als ferner API-Aufruf abläuft, bleibt die Berechnung unsichtbar, der Energieaufwand wird abstrahiert, und das Modell lässt sich weder öffnen noch anhalten oder verändern. Die Studierenden werden zu Konsumierenden eines Ergebnisses, nicht zu Lesenden eines Prozesses.

Das Seminar vertritt die gegenteilige Position: Lokale Modelle leisten bereits genug. LLaVA, auf einem Standard-Laptop betrieben, kann einen Videoframe in natürlicher Sprache beschreiben, Figuren und Gesten identifizieren, Komposition und Stimmung kommentieren. Langsamer als ein Cloud-Pendant, aber vollständig, lesbar und kostenfrei. CLIP kann ein ganzes Filmarchiv indexieren und auf eine Textabfrage hin semantisch verwandte Aufnahmen zurückliefern. Es sind produktionsnahe Arbeitsabläufe, und die Einschränkung, sie lokal laufen zu lassen, bildet selbst die didaktische Grundlage.

Die Entscheidung für den lokalen Betrieb der Modelle ist daher didaktisch und kritisch motiviert, nicht technisch.

- **Transparenz.** Das Modell ist kein Dienst, sondern ein im Hintergrund ablaufender Prozess. Was wir als Technologien bezeichnen, sind Mittel, um Ordnung in unsere Welt zu bringen (Winner, 1980). Die Studierenden sehen das Modell in Aktion. Sie können Parameter anpassen, Eingabeaufforderungen ändern und beobachten, wie sich das Ergebnis verändert.
- **Datensouveränität.** Dokumentarisches Filmmaterial, oft sensibles Material, verlässt das Atelier nicht. Für viele Bewegtbild-Projekte ist das keine akademische, sondern eine praktische Frage.
- **Keine Kosten, kontinuierliche Praxis.** Keine API-Gebühren, keine Tokens, keine Abonnements. Die Studierenden können stundenlang iterieren, ohne Budgetsorgen. Dies ist die Arbeitsweise, die künstlerische Praxis erfordert.
- **Materialbewusstsein.** Künstliche Intelligenz ist keine immaterielle Angelegenheit, sondern verkörpert sich in natürlichen Ressourcen, Energie, menschlicher Arbeitskraft und physischer Infrastruktur (Crawford, 2021). Die lokale Ausführung der Modelle macht genau dies greifbar: nicht die Abstraktion eines Cloud-Aufrufs, sondern Wärme, Lüftergeräusche und Wartezeiten als konkrete Erfahrung im Raum.

3. Lehrkonzepte

3.1 Was ist neu an dem Konzept ?

Die meisten KI-basierten Kunstkurse führen die Studierenden in kommerzielle Plattformen ein. Dieses Seminar verfolgt den umgekehrten Ansatz. Die Software wird auf den Computern der Studierenden und im Atelier installiert, wobei gänzlich auf die Cloud verzichtet wird. Diese fachspezifische Kombination aus lokaler Infrastruktur und Praxisarbeit ist im künstlerisch bzw. gestalterischen Bereich an Thüringens Universitäten beispiellos.

Das Seminar nimmt ebenso ernst, dass KI-Modelle keine neutralen Instrumente sind. Die Beschreibungen eines Frames durch LLaVA offenbaren, was das Modell zu sehen gelernt hat und was es systematisch übergeht. Die Assoziationen von CLIP innerhalb eines Bildarchivs sind geprägt von den Trainingsdaten, aus denen es hervorging. Diese Lücken und Verzerrungen sind keine Probleme, die behoben werden müssten. Sie sind der eigentliche Gegenstand des Seminars. Zentrales Erkenntnisziel ist, was ich als alternative maschinelle Erzählung verstehe: die Maschine liest dasselbe Material anders und produziert dabei eine eigenständige Perspektive, die als künstlerisches Material ernst genommen werden kann.

3.2 Kernwerkzeuge

Ollama ist die lokale Software, die man für lokale KI-Anwendungen benötigt. Eine quelloffene Anwendung, die große Sprach- und Bildmodelle direkt auf dem Rechner der Studierenden verwaltet, herunterlädt und ausführt. Ohne Internetverbindung, ohne Konto, ohne Datenübertragung. Ollama übernimmt das Laden quantisierter Modelle, stellt eine einfache API bereit, die sich aus Python oder dem Browser aufrufen lässt, und läuft gleichermaßen auf Standard-Laptops und Raspberry-Pi-Hardware. Sämtliche Werkzeuge in diesem Seminar bauen auf Ollama auf.

LLaVA (Large Language and Vision Assistant) ist ein multimodales Bild-Sprach-Modell. Es bietet Funktionen zur visuellen Beantwortung von Fragen sowie zur Bildunterschrift. Anhand eines Bildes oder eines Videobildes generiert es Beschreibungen in natürlicher Sprache: Szenenbeschreibungen, Objekterkennung, Stimmung, Komposition, Figuren und Gesten. Im Gegensatz zu Klassifikatoren mit einem geschlossenen Vokabular, die Bezeichnungen aus einer festen Liste zurückgeben, behandelt LLaVA das Bild als Gesprächspartner. Für Studierende des Bewegtbilds bedeutet dies, dass jedes Einzelbild in Text übersetzt werden kann. Man kann das Modell fragen, warum es die Szene als konfrontativ beschrieben hat, was es nicht erwähnt hat, oder es bitten, nur die Beleuchtung zu beschreiben.

CLIP (Contrastive Language Image Pretraining) leistet eine andere Form der Bilderkennung. Statt zu beschriften, misst es semantische Ähnlichkeit zwischen Bild und Text. Es kann ein ganzes Filmarchiv indexieren und die Frames zurückliefern, die einer gegebenen Phrase am nächsten kommen, etwa „shot with backlight and isolation“, „hands in motion“, „crowd and anonymity“, ohne manuelle Verschlagwortung oder Metadaten. CLIP beschreibt nicht, was es sieht. Es verortet Bilder in einem semantischen Raum, geformt durch die Millionen Bild-Text-Paare, auf denen es trainiert wurde. Meine eigene Forschung zu dem, was ich statistische Indexikalität genannt habe, also das

probabilistische Verhältnis zwischen Modell-Ausgabe und der in Trainingsdaten kodierten sozialen Welt, bildet das theoretische Rückgrat von Block 1. Wenn CLIP für dieselbe Abfrage andere Bilder zurückliefert als ein menschlicher Editor, ist diese Differenz kein Fehler. Sie ist die Weltsicht des Modells, lesbar gemacht.

Alle drei Werkzeuge sind kostenlos, quelloffen und arbeiten ohne Internetzugang. LLaVA und CLIP laufen mit praktikabler Geschwindigkeit auf Standard-Laptops mit 16 GB RAM. Für größere Batch-Aufgaben, etwa die Indexierung eines ganzen Semester-Filmarchivs, ist ein eigener Seminar-Arbeitsplatz im Projektbudget vorgesehen, der nach Projektende dauerhaft bei der Professur verbleibt.

3.3 Seminarstruktur

Das Seminar erstreckt sich über zwei Semester. Im Wintersemester (Oktober 2026 bis Februar 2027) erlernen die Studierenden die Werkzeuge, setzen sich mit der Theorie des maschinellen Sehens auseinander und besuchen ein Rechenzentrum. Im Sommersemester (April bis Juli 2027) wenden sie die Werkzeuge auf eigenes und geteiltes Archivmaterial an und entwickeln eigene gestalterische Arbeiten.

Block 1. Installation und erster Blick. „Was sieht es?“

Der erste Block des Wintersemesters verbindet praktische Installation mit theoretischer Einordnung. Beides gehört untrennbar zusammen: wer versteht, wie ein Modell trainiert wurde, liest seine Ausgaben anders. Und wer ein Modell selbst befragt hat, versteht die Theorie nicht mehr abstrakt. Alle Studierenden installieren Ollama lokal und arbeiten von Beginn an mit allen drei Werkzeugen: LLaVA beschreibt Frames in Sprache, CLIP ordnet Bilder nach semantischer Ähnlichkeit, und beide zusammen erschließen dasselbe Material auf Weisen, die sich grundlegend von menschlicher Wahrnehmung unterscheiden. Parallel dazu wird die Funktionsweise dieser Systeme im Seminar erarbeitet: Was sind Trainingsdaten? Was ist ein semantischer Raum? Woher kommt das, was das Modell „sieht“?

Block 2. Das operative Bild. „Wie versteht die Maschine?“

Ausgehend von Harun Farockis Begriff des operativen Bildes (Farocki, 2004) untersuchen die Studierenden, wie Bilderkennungsmodelle bereits in Produktionsprozesse eingebunden sind: Autotagging, Schnittvorschläge, Inhaltsmoderation (Distelmeyer, 2023). Im Laufe des ersten Semesters besuchen die Studierenden als Feldforschung ein Rechenzentrum. Die Studierenden kommen vorbereitet: Sie bringen die kritischen Fragen und Erfahrungen des Semesters mit — wie die Werkzeuge auf ihrem eigenen Rechner reagiert haben, was das Modell übersehen hat, wie lange Inferenz dauert und warum. Diese konkrete Erfahrung trifft im Rechenzentrum auf den industriellen Maßstab und macht den Größenunterschied nicht nur abstrakt begreifbar, sondern physisch erlebbar. Im Wintersemester wird eine externe Fachperson für einen Vortrag und Workshop zur kritischen KI-Praxis eingeladen.

Block 3. Das durchsuchbare Archiv. „Nach Stimmung suchen“

Im zweiten Semester wird die Arbeit mit CLIP vertieft. Die Studierenden bauen einen semantischen Suchindex über ein gemeinsames Filmarchiv auf. Aufgabe: die Suche nach einem Gefühl, einer Atmosphäre, einer Abstraktion. Was liefert das Modell? Was übersieht es? Ein gemeinsames Korpus wird parallel von Menschen und von CLIP geordnet, die Ergebnisse verglichen und diskutiert. (Steyerl, 2023)

Block 4. Alternative Erzählungen. „Was erzählt die Maschine?“

Dasselbe Archivmaterial, von der Maschine gelesen, ergibt eine andere Geschichte. LLaVA beschreibt einen Frame anders als jede Studentin, jeder Student es tun würde. CLIP ordnet ein Korpus nach Zusammenhängen, die für Menschen unsichtbar oder befremdlich sind. In diesem Block arbeiten die Studierenden gezielt mit dieser Differenz: sie montieren ihr Material entlang der maschinellen Lesart und stellen das Ergebnis der eigenen Perspektive gegenüber. Was wird sichtbar, wenn die Maschine die Erzählung führt? Was geht verloren? Die Abschlussarbeiten nehmen hier ihren Ausgangspunkt.

3.4 Abschlussarbeiten

Jede und jeder Studierende entwickelt eine Videoarbeit oder Installation, in der LLaVA oder CLIP als gestalterisches Element eingesetzt wird. KI fungiert dabei nicht als Hilfsmittel, sondern als Gegenüber mit einer eigenen Perspektive auf das Material. Begleitend entsteht ein Reflexionstext von zwei bis drei Seiten zur eigenen Arbeitspraxis. Die Abschlussarbeiten werden im Juli 2027 in einer öffentlichen Ausstellung präsentiert, als Abschluss des Seminars und als Einladung an die Hochschule, sich mit den Ergebnissen auseinanderzusetzen.

4. Zielgruppe und curricularer Kontext

Primäre Zielgruppe sind Bachelor- und Masterstudierende der Fakultät Kunst und Gestaltung der Bauhaus-Universität Weimar. Das Seminar setzt keine Vorkenntnisse in Programmierung oder Künstlicher Intelligenz voraus. Das Seminar ist als Wahlfach innerhalb der Visuellen Kommunikation konzipiert und steht ebenso Studierenden aus angrenzenden Bereichen offen (Medienkunst und Mediengestaltung, Freie Kunst oder Produktdesign). Geplante Teilnehmerzahl: 12 bis 15 Studierende pro Kurs. Nach erfolgreicher Evaluation ist die Integration in das Curriculum als empfohlenes Wahlfach vorgesehen. Die erarbeiteten Materialien werden unter CC BY 4.0 veröffentlicht und stehen damit auch anderen Einrichtungen zur Nachnutzung offen.

5. Nachhaltigkeit und Transfer

5.1 Ein wiederholbarer Prozess

Das Seminar ist als wiederholbares jährliches Format konzipiert, nicht als einmaliges Projekt. Die Vier-Block-Struktur (Werkzeuge, operatives Bild, Archivsuche, alternative Erzählungen) fügt sich in den Semesterrhythmus ein und kann nach dem Aufbau des Curriculums mit minimalem Zusatzaufwand erneut durchgeführt werden. Jede Iteration erzeugt neue studentische Arbeiten und neue Evaluationsdaten, die in den nächsten Durchgang einfließen. Das Ziel ist nicht ein Projekt, das im August 2027 endet, sondern ein Seminar, das die Professur in Eigenverantwortung fortführen kann.

5.2 Wiederverwendbare Literatur, Code und Materialien

Alle Seminarmaterialien, Installationsanleitungen, Python-Übungen, kommentierte Leseliste, Bewertungsrubrik, werden semesterbegleitend dokumentiert und auf edu-sharing Thüringen unter CC BY 4.0 veröffentlicht. Jede Einrichtung mit einem Standard-Laptop kann Block 1 innerhalb einer Stunde durchführen.

5.3 Wiederverwendbare Hardware: der Raspberry-Pi-Pool

Die im Rahmen dieses Projekts angeschafften Raspberry Pi 5-Geräte und der GPU-Computer verbleiben nach Ablauf des Stipendiums mit der darauf installierten Software dauerhaft im Lehrstuhl für Crossmediale Bewegtbilder. Jede Einheit führt Ollama und LLaVA lokal bei reduzierter Geschwindigkeit aus, ausreichend für Demonstrationen, individuelle Übungen und Workshops außerhalb des Campus. Der Pool erfüllt drei Funktionen: als Leihset für Studierende ohne geeignete Laptops, als didaktische Demonstration dessen, wie ein eingeschränktes lokales Modell aussieht, und als umverteilbare Ressource für künftige Seminare, Workshops und Outreach-Veranstaltungen an Partnerinstitutionen. Einmal angeschaffte Hardware lehrt über Jahre hinweg.

5.4 Mit KI zu arbeiten lehren. Eine übertragbare Didaktik

Über die Materialien hinaus erzeugt das Projekt einen dokumentierten didaktischen Ansatz. Wie werden KI-Bilderkennungswerkzeuge Studierenden ohne Programmierkenntnisse vermittelt? Wie lassen sich Begegnungen mit Modellverzerrungen und Modellfehlern als kreative Gelegenheiten statt als Hindernisse gestalten? Wie wird eine Exkursion so entworfen, dass sie das Denken der Studierenden verändert und nicht bloß beeindruckt? Diese methodologische Dokumentation, verfasst als kurzer praxisbasierter Beitrag für das Hochschulforum Digitalisierung, soll für alle Lehrenden in gestalterischen Studiengängen nützlich sein, die KI kritisch und zugleich praktisch unterrichten möchten.

6. Evaluation und Risiken

Der Erfolg des Seminars wird anhand von drei Kriterien beurteilt. Erstens füllen die Studierenden zu Beginn und am Ende des Seminars einen Selbsteinschätzungsbogen aus, der ihre Einschätzung eigener Kompetenzen im Umgang mit KI-Werkzeugen und deren kritische Einordnung erhebt. Zweitens wird die Qualität der Abschlussarbeiten anhand einer Bewertungsrubrik erfasst, die sowohl technische als auch gestalterische und reflexive Anteile berücksichtigt. Drittens gilt als langfristiger Indikator, ob die Professur das Seminar im Folgejahr ohne externes Budget eigenständig wiederholen kann.

Nach jedem Block wird ein kurzes formatives Feedback erhoben. Die Ergebnisse fließen direkt in die Gestaltung des nächsten Blocks ein. Diese Überarbeitung ist bewusst im Zeitplan verankert: die vorlesungsfreie Zeit im März dient explizit der Überarbeitung.

Das größte Risiko liegt in der großen Vielfalt der in den Laptops der Studierenden verwendeten Hardware. Nicht alle Geräte arbeiten bei der lokalen Anwendungen mit derselben Geschwindigkeit oder derselben Stabilität. Der Raspberry Pi und der Computer stehen als Ausweidlösung zur Verfügung und dienen zudem als Lehrmittel: Ein langsames Modell, das auf Hardware mit geringen Spezifikationen läuft, macht die Rechenarbeit besonders anschaulich.

7. Expertise und Einbindung

Das Vorhaben verbindet eigene, aktuelle Forschung unmittelbar mit der Lehre. Die folgenden Arbeiten bilden seine Grundlage:

- Essay zur Materialität von KI-Hardware und der politischen Ökonomie der Bildgenerierung (opscience.ub.uni-mainz.de, 2024)
- Arbeit zur Stereotypen-Kodierung in Vision-Language-Modellen und zum Konzept der statistischen Indexikalität (Academia.edu, 2024)
- Vortrag zu Tokensprachen und politischer Repräsentation in großen Sprachmodellen (Chaos Communication Congress 2025)
- „Steering Through the Inner Residue“, Videoinstallation mit Leon-Etienne Kühr (Museum Angewandte Kunst Frankfurt, 2026)
- Bereits erprobte Lehrveranstaltungen zu Creative Coding und lokaler KI an der Bauhaus-Universität Weimar

Das Vorhaben ist in der Professur Crossmediales Bewegtbild verankert und wird von Prof. Jakob Hüfner unterstützt. Seine langjährige Praxis mit experimentellen Methoden im Bewegtbild schafft einen produktiven Rahmen für das Seminar: Die Auseinandersetzung mit KI-Werkzeugen tritt hier nicht in Konkurrenz zur konventionellen Filmarbeit, sondern in Dialog mit ihr. Studierende, die in einer Professur mit dieser künstlerischen Haltung ausgebildet werden, bringen bereits eine Sensibilität für das Bild mit, die die kritische Arbeit mit maschineller Wahrnehmung erst fruchtbar macht.

8. Fellows-Austausch

Wie bewertet man kritische Reflexionskompetenz in gestalterischen Studiengängen? Wie unterscheidet sich eine gelungene künstlerische Auseinandersetzung mit KI von einer bloßen Anwendung? Für diese Fragen gibt es innerhalb der eigenen Professur wenig Vergleichsmöglichkeit. Der Austausch mit Fellows aus anderen Disziplinen und Hochschulen bietet die Chance, Evaluationsformate zu vergleichen, didaktische Erfahrungen zu teilen und das eigene Vorhaben von außen prüfen zu lassen.

9. Literaturverzeichnis

- Crawford, K. (2021). *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press.
- Distelmeyer, J. (2023). Which Operativity? On Political Aspects of Operational Images and Sounds. *Interface Critique*, vol. 4. Auch: harun-farocki-institut.org
- Farocki, H. (2004). *Phantom Images*. Public, no. 29, pp. 12–22.
- Hunger, F. (2023). Unhype Artificial 'Intelligence'! A proposal to replace the deceiving terminology of AI. Zenodo / Training the Archive, HMKV Dortmund. <https://zenodo.org/records/7524493>
- Liu, T.-C. (2024a). On the Materiality of Artificial Intelligence, Article Contribute to un/learn ai : integrating AI in aesthetic practices, Volume 3. pg.218-223. Published by Forschungsprojekt KITeGG.
-
- Liu, T.-C. & Kühr, L.-E. (2024b). *Stereotype Encoding: How AI Images learn Cultural Bias*. https://www.academia.edu/165438573/Stereotype_Encoding_How_AI_Images_Learn_Cultural_Bias
- Liu, T.-C. & Kühr, L.-E. (2026). *Steering Through the Inner Residue*. Skript zur Videoinstallation. Museum Angewandte Kunst, Frankfurt.
- Paglen, T. (2016). *Invisible Images. Your Pictures Are Looking at You*. *The New Inquiry*, 8 December.
- Penny, S. (1995). Consumer Culture and the Technological Imperative. In: *Critical Issues in Electronic Media*, pp. 47–73. SUNY Press.
- Steyerl, H. (2023). Mean Images. *New Left Review*, 140/141, pp. 82–97.
- Winner, L. (1980). Do Artifacts Have Politics? *Daedalus*, vol. 109, no. 1, pp. 121–136.