

André Weßel/Friederike von Gross (Hrsg.)

Aufwachsen mit KI

**Medienbildung, Datenschutz und digitale Selbstbestimmung
im Kindes- und Jugendalter**

André Weßel/Friederike von Gross (Hrsg.)

Aufwachsen mit KI

**Medienbildung, Datenschutz und digitale Selbstbestimmung
im Kindes- und Jugendalter**

Beiträge aus Forschung und Praxis

Prämierte Medienprojekte

André Weßel/Friederike von Gross (Hrsg.)

Aufwachsen mit KI

**Medienbildung, Datenschutz und digitale Selbstbestimmung
im Kindes- und Jugendalter**

Beiträge aus Forschung und Praxis – Prämierte Medienprojekte

Dieser Band wurde gefördert vom

Bundesministerium für Bildung, Familie, Senioren, Frauen und Jugend (BMBFSFJ)

Herausgeber

Gesellschaft für Medienpädagogik und Kommunikationskultur
in der Bundesrepublik Deutschland e.V. (GMK)

Anschrift

GMK-Geschäftsstelle
Oberrnstr. 24 a
33602 Bielefeld
fon 0521/677 88
email gmk@medienpaed.de
homepage www.gmk-net.de

Redaktion

André Weßel
Dr. Friederike von Gross
Tanja Kalwar

Lektorat

Tanja Kalwar

Titellillustration

kopaed

Druck

Memminger MedienCentrum, Memmingen

© kopaed 2026

Arnulfstraße 205
80634 München

fon 089/688 900 98
fax 089/689 19 12
email info@kopaed.de
homepage www.kopaed.de

ISBN 978-3-96848-210-1

Frank Schlegel

Was tun gegen Deepfakes?

Und worum es eigentlich geht!

Zu Beginn der technologischen Disruption, so Anfang 2023, als wir mit KI-Bildgeneratoren wie Midjourney herumspielten, stand eine Prognose im Raum: Bald kommen die KI-Videos und mit ihnen die KI-Bilderflut! Und: Wir werden KI-Fakes von echten Inhalten nicht mehr unterscheiden können. Mitte 2025 erschien dann Googles Videogenerator Veo 3. Das war der Dammbruch. Gängige Plattformen wurden mit *AI-Slop* (KI-Müll) überschwemmt. Das ging sehr schnell. Und mit den Katzenvideos kamen die Deepfakes, die auf TikTok munter viral gingen.

Gefakte Werbebotschaften mit renommierten Ärzt*innen, Fake-Zitate von Politiker*innen oder komplette Fake-Profile bekannter Persönlichkeiten – obwohl solche Inhalte gemeldet werden, will TikTok hier keine Verstöße gegen seine Community-Richtlinien erkennen. Die Videos verschwinden oft erst nach massiven journalistischen Recherchen oder einem erfolgreichen Rechtsstreit gegen TikTok (Nicolaus/Scherndl 2025; Nicolaus/Thom 2026).

Konjunktur für Desinformation

Diese Deepfakes verbreiten sich in einer krisen-geprägten Zeit, in der wir vielerorts lesen und hören, dass „Verschwörungstheorien“ „Konjunktur“ haben (Deutschlandfunk 2026). Im Fahrwasser der Epstein-Files übertrumpfen sich Verschwörungserzählungen in ihrer Absurdität. Anything goes, auch embryonale Gewürzmittel, die von Kannibal*innen in unsere Kartoffelchips gepanscht wurden (Hecht 2026). Vor jeder politischen Wahl wird eine demokratiegefährdende Desinformationswelle befürchtet. Der im Rahmen des Weltwirtschaftsforums entstandene Global Risks Report 2026 stuft „Misinformation and disinformation“ als derzeit weithöchste globale Gefahr ein (World Economic Forum 2026).

Kein Wunder, dass auch Lehrkräfte dieses Themenfeld als dringlich empfinden. Das merke ich an den hohen Teilnehmendenzahlen, wenn ich Workshops wie „Deepfakes: Desinformation in Zeiten Künstlicher Intelligenz“ auf einem Pädagogischen Fachtag oder einer Tagung anbiete. Ich möchte im Folgenden durchspielen, wie ich solche Angebote derzeit gestalte, welche Erfahrungen ich dabei gemacht habe und welche Fragen sich mir stellen.

Übrigens gestalte ich Workshops für Jugendliche in ihrer Grundstruktur genauso wie für Erwachsene, weil generative KI für alle gleich neu ist. Ich konzentriere mich hier auf meine Erfahrungen in der Erwachsenenfortbildung, weil ich bei erwachsenen Zielgruppen eine zunehmende Desorientierung zu diesem Thema wahrnehme. Lehrkräfte und Akteur*innen in der außerschulischen Bildung müssen sich aber sicher fühlen, um wiederum Jugendlichen Orientierung bieten zu können.

Wie können wir Deepfakes erkennen?

Ich starte den Workshop mit dem Spiel „KI oder nicht“? Dabei werden die Teilnehmenden mit Bildern und Videos konfrontiert und müssen entscheiden, ob sie es mit einem authentischen oder einem KI-generierten Medium zu tun haben. Dafür habe ich zuerst ein *Kahoot!* gebastelt, inzwischen nutze ich das Online-Spiel des KI-Detektors sightengine. Bei Kindern wie Erwachsenen ist die Methode sehr beliebt. Wenn es schnell gehen muss, spiele ich mit dem Plenum zusammen über den Beamer und wir entscheiden per Mehrheitsvotum, ob wir ein Bild als echt oder fake einstufen. Auf der „24. Bundesweiten Expertinnen- und Expertentagung Lehrkräfteausbildung“ habe ich Lehrkräfte in Kleingruppen eingeteilt und sie mithilfe des

Play with AI Videos

Take a look, try to recognize AI videos.



Abb.1: Auf sightengine.com können wir testen, ob wir KI-Videos erkennen können.

Spiels eine Checkliste erstellen lassen, wie man KI-generierte Inhalte erkennen kann. In der anschließenden Diskussion gaben die Lehrkräfte allerdings an, dass das inzwischen nicht mehr möglich sei. Diese Erkenntnis ist Ziel der Übung. Ich halte solche Checklisten, die sich an technischen Indizien abarbeiten, für wenig hilfreich, weil sie schnell inaktuell werden. Die klicksafe-Checkliste „Deepfake Detectives“ führt zum Beispiel auf: „Ohren, Nase oder Zähne haben seltsame Formen“ (klicksafer o. J.). Verzernte Gesichter, verdrehte Buchstaben und Hände mit sechs Fingern – im Fall von KI-Bildern sind das Artefakte aus der KI-Steinzeit, die größtenteils seit Midjourneys Version 6 Update (2024) und endgültig mit Googles Modell Nano Banana (2025) verschwunden sind.

KI-Detektoren – besser als das menschliche Auge?

In der Regel fragt an der Stelle jemand, ob KI-Detektoren helfen können. Also zeige ich, wie ich ein Bild von acht kostenfrei nutzbaren KI-Detektoren untersuchen lasse. Die Einschätzungen der Modelle variieren meist stark. Es fällt auf: Sind Detektoren nicht aktuell, werden sie schnell nutzlos.



Abb. 2: Hier hatte Google Gemini den Auftrag, ein Selfie im Stil eines KI-Bilds aus dem Jahr 2023 darzustellen.

Warum sind Deepfake-Generatoren gegen KI-Detektoren generell im Vorteil? Sie können leicht trainiert werden: Ein Modell, das Deepfakes produzieren soll, kann von Modellen, die Deepfakes erkennen sollen, lernen. Es versucht, die Detektoren so lange zu überlisten, bis es schließlich gelingt. Und schon sind die Detektoren veraltet (Jensen 2024).

Außerdem demonstriere ich, wie ich mithilfe einfacher Prompts einen KI-Detektor überlisten kann. Zunächst listet ChatGPT auf, mit welchen Mitteln ein Bild besonders realistisch wirken kann. Daraus erstellt ChatGPT dann direkt einen Prompt für einen Bildgenerator. Mit einem so erstellten Bild konnte ich einen KI-Detektor in die Irre führen. Sowieso gilt: Wo authentische Bilder mit KI nachbearbeitet werden, verschwimmen die Muster.

Die gute, alte Informationskompetenz

Meine Workshop-Teilnehmenden haben nun erfahren, dass sie Deepfakes nicht zuverlässig als solche identifizieren können. Diese Erkenntnis ist Ausgangspunkt unseres weiteren Nachdenkens darüber, wie wir Desinformation im Netz begegnen können. Also, wie gehen

wir damit um? Die Antwort ist nicht neu: mit Informationskompetenz. Man muss halt die Quellen checken, im Zweifel gegenrecherchieren, kritisch sein. Manipulative Muster erkennen können. Tatsächlich machen die Entwicklungen des KI-Zeitalters die Förderung von Informationskompetenz noch dringlicher, weil manipulierte Inhalte stark zunehmen.

Lehrkräfte offenbaren mir im Gespräch allerdings oft, dass sie sich eine Beurteilung aktueller Informationen nicht zutrauen. „Zu den Epstein-Files kann ich auch nichts sagen.“ „Man weiß gar nicht mehr, wem man noch trauen kann.“ Das ist ein großes Problem, denn allgemeine Verunsicherung im Angesicht einer diffusen Nachrichtenlage ist ein Etappenziel von Akteur*innen, die Desinformation bewusst in Umlauf bringen.

Wenn Jugendliche erkennen, dass sie im Netz mit unseriösen Infos konfrontiert werden und sich hilfesuchend an eine Lehrperson wenden, sollte ihnen Orientierung geboten werden.

In Workshops mit jungen Erwachsenen gehe ich deshalb so vor: Ich bitte die Lernenden, eigene Postings und Meldungen mitzubringen, deren Wahrheitsgehalt sie nicht einschätzen können. Oder ich stelle zu Beginn des Workshops die Frage „Was ist gestern alles im Internet passiert?“ Wenn ich mit Inhalten konfrontiert werde, die mir unbekannt sind, mache ich deutlich, dass ich dazu noch keine Position habe, aber weiß, wie ich mir eine bilden kann. Und dann geht es in den gemeinsamen Rechercheprozess. In diesem Prozess soll deutlich werden: Es gibt überprüfbare Informationen und solche, die nur behauptet werden. Es gibt auch mehrdeutige Informationen. Und es gibt Meinungen.

Ich möchte Lehrkräfte befähigen, solche Debunking-Prozesse durchzuführen und sich den Rollenwechsel zuzutrauen, gemeinsam mit ihren Schüler*innen nach Antworten zu suchen. Also versorge ich Lehrkräfte mit entsprechendem Handwerkszeug: einer Übersicht zu den gängigen Faktencheck-Seiten wie Correctiv, Mimikama und Faktenfuchs, Methoden wie der Bildersuche Rückwärts und Projekten wie *So geht Medien!*, *#DigitalCheckNRW* und *klicksafe*.



Abb. 3: Ein bisschen KI-generierte Überbelichtung – und schon konnte dieses Bild den KI-Generator überlisten.

Um zu beweisen, dass sie die notwendige Informationskompetenz schon mitbrächten, ließ ich im Herbst 2025 das Kollegium eines Berufskollegs gemeinsam ein konkretes Desinformations-Beispiel nachrecherchieren. Und das ging nach hinten los...

Ein Beispiel aus dem Medien-Alltag: Der „Held von Hamburg“

In meinem Beispiel ging es um den „Helden von Hamburg“. In sozialen Netzwerken wurde behauptet:

„Der Syrer, der laut Medienberichten am 23. Mai 2025 eine Angreiferin am Hauptbahnhof in Hamburg aufgehalten hat, sei ein KI-Fake. Zwei Bilder würden ihn in derselben Pose an verschiedenen Orten zeigen, einmal an einem Bahnsteig und einmal vor dem Brandenburger Tor in Berlin.“ (Kutzner 2025)

Der Correctiv-Faktencheck stellte richtig:

„Falsch. Den Mann gibt es, er hat mit mehreren Journalisten über den Moment am Bahnhof in Hamburg gesprochen. Das Bild, das den jungen



Abb.4: Der „Held von Hamburg“: Screenshots von Postings auf X und Facebook

Mann an einem Bahnsteig zeigt und in Sozialen Netzwerken verbreitet wird, ist KI-generiert – aus Medienberichten stammt es nicht. Das Foto, das ihn vor dem Brandenburger Tor zeigt, ist dagegen echt.“ (ebd.)

Die Lehrkräfte bekamen als Grundlage lediglich den Screenshot einer Meldung des SPIEGELS sowie Screenshots von Postings auf X und Facebook (siehe Abb.4). Die Desinformationsstrategie ist hier besonders perfide, da sie sich auf die „Lügner*innendividende“ stützt: Der durch Deepfakes verursachte Vertrauensverlust in den Wahrheitsgehalt audiovisueller Medien kann genutzt werden, um authentische Bildquellen mit dem Deepfake-Vorwurf infrage zu stellen (Pawelec 2024).

Mit der Frage „Was ist wirklich passiert?“ entließ ich die Lehrkräfte in die Recherchephase. Hier war alles erlaubt, auch die Nutzung von KI-Chatbots wie ChatGPT und KI-Suchmaschinen wie Perplexity. Im Anschluss an die Recherche sollte das Plenum sammeln, welche Recherchestrategien sich als zielführend erwiesen.

„Die Öffentlich-Rechtlichen? Da stimmt auch nicht immer alles.“

Das Ergebnis: Rund zwei Drittel hatten den Fall erfolgreich gegenrecherchiert. Die erfolgreichen Gruppen nannten als Quellen Correctiv und öffentlich-rechtliche Berichte von NDR und Tagesschau. Einige Lehrkräfte kritisierten jedoch, dass die Medien ein KI-Bild verwendet hätten, obwohl der „Held von Hamburg“ doch real sei. Einige hielten die komplette Meldung für eine Lüge und äußerten, die Inhalte der Öffentlich-Rechtlichen würden „auch nicht immer“ stimmen. Außerdem hatte ein Lehrer „irgendwo gelesen“, dass Correctiv linkspolitisch gefärbt sei.

Das war neu für mich. Aus der Fortbildung von Multiplikator*innen bin ich die Frage gewohnt, wie man denn mit Jugendlichen umgehen könne, die ein radikales Weltbild vertreten oder journalistischen Quellen schlicht nicht vertrauen. Hier hatte ich nun Lehrkräfte vor mir, die ihren Google-Funden mehr Glauben schenken als einer dpa-Meldung.

Was, wenn ein Beweis nicht als Beweis, eine Quelle nicht als Quelle akzeptiert wird? Ich muss das Medienvertrauen ansprechen. Die Upgrade-Democracy-Studie „Verunsicherte Öffentlichkeit“ untersucht Medienvertrauen als Schlüsselkategorie im Umgang mit Desinformation. Denn:

„Befragte mit geringem Medienvertrauen haben ein weiteres Begriffsverständnis von Desinformationen. Sie fassen häufiger auch unbeabsichtigte Falschmeldungen darunter und sind häufiger der Meinung, es gehe dabei hauptsächlich darum, alternative Meinungen zu diskreditieren. Ferner nehmen sie häufiger Desinformationen wahr (fast die Hälfte), halten inländische Akteur:innen, Politiker:innen, Journalist:innen und auch die Bundesregierung eher für die Verursacher:innen und vermuten als Motiv überdurchschnittlich oft, dass damit von Skandalen und politischer Unfähigkeit abgelenkt werden soll.“ (Bernhard et al. 2024)

Who watches the Watchmen?

Den individuellen Ursachen geringen Medienvertrauens kann ich im Plenum nicht auf den Grund gehen. Wie aber kann ich akut reagieren, wenn Faktenchecks und Öffentlich-Rechtliche infrage gestellt werden? Ich versuche, ein Verständnis dafür zu schaffen, unter welchen Bedingungen etwa ein Correctiv-Artikel überhaupt entstanden ist. Da wir Absender vermeintlicher Desinformation checken sollten, überprüfen wir auch Correctiv auf Herz und Nieren. Es wird deutlich: Als gemeinnützige Organisation legt Correctiv ihre Zuwendungen detailliert offen, es gibt ein umfangreiches Redaktionsstatut und einen Aufsichtsrat. Damit wird man die gefestigte Meinung, bei Correctiv handele es sich um ein links-grünes Meinungsorgan, wahrscheinlich nicht ins Wanken bringen. Aber man kann zeigen, dass Correctiv nicht im Diffusen arbeitet, sondern in einen nachvollziehbaren Regelrahmen eingebettet ist (CORRECTIV o. J.).

Auch für das Regelsystem öffentlich-rechtlicher Medien versuche ich Verständnis zu schaffen. Wer kontrolliert die? Wie ist ihr Verhältnis zum Staat geregelt? Was ist ein Pressestatut? Was passiert, wenn eine Falschnachricht veröffentlicht wird? Wie funktioniert eigentlich das duale Rundfunksystem? Im Workshop zeigte sich, dass vielen Lehrkräften Basiswissen zu Medienorganisationen und der Arbeitsweise von Journalist*innen fehlte. Dieses Basiswissen ist aber erforderlich, um journalistische Medien fundiert kritisieren zu können. Journalismusforscher*innen fordern deshalb die Förderung von „Journalismuskompetenz als spezifizierte Medienkompetenz“ (Beiler/Krüger 2023). Für jüngere Zielgruppen gibt es verschiedene Methoden, um journalistische Arbeit erfahrbar zu machen, zum Beispiel das Lernspiel *Nachrichtenschmacher* von Planet Schule (Schmid 2025). Wir müssen Journalismuskompetenz dringend auch in der Erwachsenenbildung fördern.

Medienkritik? Unbedingt!

Auch ein kritisch-reflektierter Umgang mit Medien ist eine Säule von Journalismuskompetenz. Ich mache meinen Lerngruppen deutlich: Kritik etwa an öffentlich-rechtlichen Medien, die nicht auf einem pauschalen Lügenpresse-Vorwurf, sondern auf nachvollziehbaren Argumenten fußt, ist erlaubt und erwünscht. Sie ist Teil des demokratischen Aushandlungsprozesses.

Medienkritik findet auch in der Medienlandschaft selbst statt. Das konnte ich im Februar 2026 anhand eines aktuellen Falls skizzieren und dadurch einen Vorwurf an die Medien umformulieren zu einem Beispiel für die Selbstkorrekturfähigkeit des Mediensystems:

Am 15. Februar wurde im *heute journal* in einem Beitrag über die US-Migrationsbehörde ICE ein KI-generiertes Video gesendet. Obwohl in dem Video das Logo von OpenAI Sora zu sehen war, gab das ZDF es als echt aus. Es folgte ein Sturm der Kritik in verschiedenen Medien, auf die der Sender erst nach und nach reagierte. In meinem Workshop am 27. Februar äußerten Lehrkräfte, so ein Fauxpas würde das Vertrauen in Medien endgültig zerstören. Ich lud dazu ein, schnell mittels Google und der KI-Suchmaschine Perplexity nachzuschlagen, welche Konsequenzen das sogenannte ZDF-KI-Gate bereits nach sich gezogen hatte. Es zeigte sich: Ein breites Medienecho hatte auch die Aufarbeitung des Fehlers journalistisch begleitet, das ZDF hatte den Fehler richtiggestellt, sich entschuldigt und die verantwortliche Korrespondentin abberufen (Rosenkranz/Graf 2026). Das demonstriert zum einen die öffentliche Selbstkontrolle des Journalismus insgesamt und zum anderen verankerte Regeln wie die Korrekturpflicht. Außerdem konnte ich anhand des Falls die im Englischen zutreffende Unterscheidung zwischen *misinformation*, der unbeabsichtigten Falschinformation, und *disinformation*, also der gezielt verbreiteten Täuschung, deutlich machen.

In Fortbildungen finde ich mich häufig in der Rolle des ÖRR-Verteidigers wieder, weil ich für Debunking journalistische Quellen benötige. Ich mache aber durchaus deutlich, dass ich nicht blind vertraue, sondern auch Kritikpunkte am ÖRR habe. Zum Beispiel werden allgemein Femizide in den Nachrichten unterproportional und Straftaten durch Ausländer überproportional oft genannt. Strukturell bedingte Gewalttaten werden dabei meist als Einzelfälle dargestellt (Meltzer 2022). Dieses (egal ob bewusste oder unbewusste) Agenda-Setting formt die Wahrnehmung der Öffentlichkeit und kann letztendlich ähnliche Auswirkungen hervorrufen wie intentional verbreitete Desinformation.

Noch mehr Elefanten im Raum

Derzeit sind Deepfakes ein Anlass, um Fortbildungsangebote im Themenfeld Desinformation anzubieten. Diese Angebote sollten ein Türöffner sein für eine tiefe Beschäftigung mit Informations- und Journalismuskompetenz. Beim Üben von Debunking-Methoden rücken auch KI-Chatbots und KI-Agenten in den Fokus, weil diese neben Kurzvideos selbstverständlich als alltägliche Informationsquelle genutzt werden.

Ich habe im Beitrag beschrieben, dass ich in Workshop-Situationen damit konfrontiert werde, die von mir präsentierten Faktencheck-Quellen wären linkspolitisch verzerrt. Da fehlt das Medienvertrauen für Debunking. Und präventive Strategien wie Prebunking zielen darauf, langfristig eine Resilienz gegenüber Beeinflussungsversuchen zu entwickeln (IDZ et al. 2026). In der direkten Konfrontation benötige ich aber eine Interventionsstrategie. Hier muss ich mich als Medienpädagoge weiterbilden.

Natürlich beschäftige ich mich auch mit den Ursachen für geringes Medienvertrauen. Eine Auffälligkeit: Von den in der Upgrade-Democracy-Studie Befragten mit niedrigem Medienvertrauen beabsichtigen 58 Prozent, die AfD zu wählen – 17 Prozent die CDU/CSU, 6 Prozent die SPD und 5 Prozent die Linken (Bernhard et al. 2024: 13). Das verwundert

nicht. Während AfD-Politiker*innen journalistische Medien als „Lügenpresse“ oder „Systemmedien“ diffamieren, hat die AfD in sozialen Netzwerken ein „Ökosystem“ für Desinformation etabliert (Amadeu Antonio Stiftung 2024). Zivilgesellschaftliche Organisationen und medienpädagogische Projekte, die Desinformation bekämpfen, sind der AfD folglich ein Dorn im Auge. In einer Kleinen Anfrage an den Bundestag äußerte die AfD-Fraktion den Verdacht, das Thema Desinformation diene als Vorwand, „um unbequeme Meinungsäußerungen gezielt zu delegitimieren.“ Sie fragte weiter, „welche Maßnahmen die Bundesregierung zur Vermeidung möglicher Verzerrungen oder Voreingenommenheit (Bias) in den Beurteilungen und Ergebnissen von beteiligten zivilgesellschaftlichen Akteuren“ trafe (Deutscher Bundestag 2025). Das wirkt tatsächlich harmlos im Vergleich zur vorangegangenen Anfrage der CDU/CSU im Februar 2025. Hier stellt die Fraktion in 551 Fragen die politische Neutralität von u. a. Correctiv und dem Bundesprogramm *Demokratie leben!* infrage, nachdem es Proteste gegen das gemeinsame Abstimmungsverhalten von Union und AfD gegeben hatte. Die Unions-Fraktion führte aus:

„Die Kritik an der Einflussnahme gemeinnütziger Organisationen geht jedoch über einzelne Proteste hinaus. Manche Stimmen sehen in den Nichtregierungsorganisationen (NGOs) eine Schattenstruktur, die mit staatlichen Geldern indirekt Politik betreibt. Laut einem Bericht der ‚Welt‘ erhalten zahlreiche NGOs, die sich öffentlich politisch links positionieren, finanzielle Mittel aus Bundesministerien. Dies stellt ein Spannungsverhältnis dar, denn wenn diese Organisationen aktiv in politische Meinungsbildung eingreifen, könnte dies ein Verstoß gegen die demokratische Grundordnung sein.“ (Deutscher Bundestag/CDU/CSU 2025)

Verlinkt wurde an dieser Stelle keine investigative Recherche, sondern ein in der Welt veröffentlichter Kommentar, der in den „angeblichen NGOs“ einen „Staat im Staate“ mit gefährlicher Macht sah. Hier begegnet uns

die Verschwörungserzählung vom „Deep State“, die von rechtsextremen Bewegungen wie QAnon verbreitet und von Donald Trump und der AfD aufgegriffen wurde (Richter 2025). Politische Konsequenzen folgten: Im März 2026 gab Bundesfamilienministerin Prien bekannt, zahlreichen *Demokratie leben!*-Projekten die Fördermittel zu entziehen.

Okay. Warum wollten wir uns nochmal mit Deepfakes und Desinformation auseinandersetzen? Und wie tut man das „politisch neutral“? Und wie finanzieren wir das in Zukunft?

Autor

Frank Schlegel: digitaldurstig.de; freiberuflicher Medientrainer; bietet Fortbildungen für Bildungsakteur*innen in unterschiedlichen Arbeitsfeldern an; seine Schwerpunkte: kreativer Einsatz von KI-Modellen in der Bildungsarbeit, Transformation von Lernen in der postdigitalen Gesellschaft, Umgang mit Desinformation und Deepfakes.

Literatur

- Amadeu Antonio Stiftung (2024): Die „Lügenpresse!“-Rufe der AfD werten Medien und Meinungsfreiheit ab. Abrufbar unter: www.amadeu-antonio-stiftung.de/die-luegenpresse-rufe-der-afd-werten-medien-und-meinungsfreiheit-ab-117097/ [Stand: 24.03.2026].
- Beiler, Markus/Krüger, Uwe (2023): Vertrauenskrise des Journalismus. Informationen zur politischen Bildung Nr. 355. Abrufbar unter: www.bpb.de/shop/zeitschriften/izpb/medienkompetenz-355/539985/vertrauenskrise-des-journalismus/ [Stand: 20.03.2026].
- Bernhard, Lukas/Schulz, Leonie/Berger, Cathleen/Unzicker, Kai (2024): Verunsicherte Öffentlichkeit. Sorgen in Deutschland und den USA wegen Desinformationen. Abrufbar unter: <https://www.bertelsmann-stiftung.de/de/publikationen/publikation/did/verunsicherte-oeffentlichkeit> [Stand: 20.03.2026].
- Correctiv (o. J.): Finanzen & Förderer. Abrufbar unter: <https://correctiv.org/ueber-uns/finanzen/> [Stand: 20.03.2026].
- Deutscher Bundestag (2025): Desinformation – Politische Zielsetzungen und Maßnahmen der Bundesregierung im digitalen Meinungskampf. Kleine Anfrage, Drucksache 21/1481. Abrufbar unter: <https://dserver.bundestag.de/btd/21/014/2101481.pdf> [Stand: 24.03.2026].
- Deutscher Bundestag (2025): Politische Neutralität staatlich geförderter Organisationen. Kleine Anfrage der Fraktion der CDU/CSU, Drucksache 20/15035. Abrufbar unter: <https://dserver.bundestag.de/btd/20/150/2015035.pdf> [Stand: 24.03.2026].
- Deutschlandfunk (2026): Illuminati, Deep State & Co. Warum sich Verschwörungstheorien so hartnäckig halten. Abrufbar unter: <https://www.deutschlandfunk.de/verschwoerungstheorien-deep-state-antisemitismus-illuminati-100.html> [Stand: 18.03.2026].
- Hecht, Janina (2026): Xavier Naidoo am Rande einer Demo in Berlin: „Kinder werden gefressen“. Abrufbar unter: <https://www.tagesschau.de/inland/regional/badenwuerttemberg/swr-xavier-naidoo-am-rande-einer-demo-in-berlin-kinder-werden-gefressen-100.html> [Stand: 18.03.2026].
- Institut für Demokratie und Zivilgesellschaft (IDZ)/Gesellschaft für Medienpädagogik und Kommunikationskultur (GMK)/toneshift (2026): Kompetenzen im Umgang mit Desinformation und KI. Atlas Wahrheit und Wissen im digitalen Wandel. Abrufbar unter: <https://atlas.machine-vs-rage.net/kompetenzen-im-umgang-mit-desinformation-ki/> [Stand: 23.03.2026].
- Jensen, Magnus Stenaa (2024): Detecting AI-manipulated content is a challenging arms race. Abrufbar unter: www.dtu.dk/english/newsarchive/2024/03/detecting-ai-manipulated-content-is-a-challenging-arms-race [Stand: 18.03.2026].
- klicksafe (o. J.): Checkliste „Deepfake Detectives“. Abrufbar unter: www.klicksafe.de/fileadmin/cms/download/Material/Päd._Praxis/Zusatzmaterial_Checkliste-Deepfake-Detectives_klicksafe.pdf [Stand: 18.03.2026].
- Kutzner, Steffen (2025): Wie mit einem KI-Bild Stimmung gegen den „Helden von Hamburg“ gemacht wird. Abrufbar unter: <https://correctiv.org/faktencheck/2025/05/28/messerangriff->

- hamburg-was-hinter-ki-bild-vom-held-von-hamburg-steckt/ [Stand: 19.03.2026].
- Meltzer, Christine (2022): Die Darstellung von Gewalt gegen Frauen in den Medien. Abrufbar unter: www.bpb.de/themen/gender-diversitaet/femizide-und-gewalt-gegen-frauen/515609/die-darstellung-von-gewalt-gegen-frauen-in-den-medien/ [Stand: 20.03.2026].
- Nicolaus, Kimberly/Scherndl, Gabriele (2025): Mel-den zwecklos: Tiktok ignoriert Desinformation zur Bundestagswahl. Abrufbar unter: <https://correctiv.org/faktencheck/hintergrund/2025/02/21/melden-zwecklos-tiktok-ignoriert-desinformation-zur-bundestagswahl/> [Stand: 18.03.2026].
- Nicolaus, Kimberly/Thom, Paulina (2026): „Massiv rufschädigend“: Tiktok ignoriert Deepfake-Werbung mit Ärzten. Abrufbar unter: <https://correctiv.org/faktencheck/hintergrund/2026/02/13/rote-bete-kapseln-tiktok-ignoriert-deepfake-werbung-mit-aerzten/> [Stand: 18.03.2026].
- Renner, Martin Erwin/Frömming, Götz/Gläser, Ronald/Helferich, Matthias/Hess, Nicole/Wendorf, Sven/Gauland, Alexander/Teich, Tobias (2025): Desinformation – Politische Zielsetzungen und Maßnahmen der Bundesregierung im digitalen Meinungskampf. Kleine Anfrage der Fraktion der AfD. Abrufbar unter: <https://dserver.bundestag.de/btd/21/014/2101481.pdf> [Stand: 24.03.2026].
- Richter, Marie Eleonore (2025): Deep State: Wie sich der Mythos ausbreitet. Faktenfuchs. Abrufbar unter: www.br.de/nachrichten/netzwelt/deep-state-wie-sich-der-mythos-ausbreitet-faktenfuchs,Ufht7iy [Stand: 24.03.2026].
- Rosenkranz, Boris/Graf, Alexander (2026): Die falschen ICE-Videos im „heute journal“ und die katastrophale Reaktion des ZDF. Abrufbar unter: <https://uebermedien.de/114229/die-falschen-ice-videos-im-heute-journal-und-die-katastrophale-reaktion-des-zdf/> [Stand: 23.03.2026].
- Pawelec, Maria (2024): Politische Manipulation und Desinformation. Abrufbar unter: www.bpb.de/lernen/bewegt-bild-und-politische-bildung/556305/politische-manipulation-und-desinformation/ [Stand: 19.03.2026].
- Schmid, Johanna (2025): Nachrichtenmacher. Lernspiel zur Medienkompetenz. Abrufbar unter: www.planet-schule.de/thema/nachrichtenmacher-lernspiel-100.html [Stand: 20.03.2026].
- World Economic Forum (2026): The Global Risks Report 2026. Geneva: World Economic Forum.

Lizenz

Der Artikel steht unter der Creative Commons Lizenz **CC BY-SA 4.0**. Der Name des Urhebers soll bei einer Weiterverwendung genannt werden. Wird das Material mit anderen Materialien zu etwas Neuem verbunden oder verschmolzen, sodass das ursprüngliche Material nicht mehr als solches erkennbar ist und die unterschiedlichen Materialien nicht mehr voneinander zu trennen sind, muss die bearbeitete Fassung bzw. das neue Werk unter derselben Lizenz wie das Original stehen. Details zur Lizenz: <https://creativecommons.org/licenses/by-sa/4.0/legalcode>.

Einzelbeiträge werden unter www.gmk-net.de/publikationen/artikel veröffentlicht.