



Simon Kruschinski, Pablo Jost,
Hannah Fecher, Tobias Scherer

Künstliche Intelligenz in politischen Kampagnen

Akzeptanz, Wahrnehmung und Wirkung

OBS-Arbeitspapier 75
ISSN: 2365-1962 (nur online)

Herausgeber:
Otto Brenner Stiftung
Can Gülcü
Wilhelm-Leuschner-Straße 79
D-60329 Frankfurt am Main
Tel.: 069-6693-2810
E-Mail: info@otto-brenner-stiftung.de
www.otto-brenner-stiftung.de

Für die Autor*innen:
Simon Kruschinski
Johannes Gutenberg-Universität Mainz
simon.kruschinski@uni-mainz.de

Redaktion & Lektorat:
Robin Koss (OBS)

Satz und Gestaltung:
Isabel Grammes, think and act

Titelbild:
Erstellt mit Midjourney

Redaktionsschluss:
13. Januar 2025

Hinweis zu den Nutzungsbedingungen:



Dieses Arbeitspapier ist unter der Creative Commons „Namensnennung – Nicht-kommerziell – Weitergabe unter gleichen Bedingungen 4.0 International“-Lizenz (CC BY-NC-SA 4.0) veröffentlicht.

Die Inhalte sowie Grafiken und Abbildungen dürfen, sofern nicht anders angegeben, in jedwedem Format oder Medium vervielfältigt und weiterverbreitet, geremixt und verändert werden, sofern keine Nutzung für kommerzielle Zwecke stattfindet. Ferner müssen angemessene Urheber- und Rechteangaben gemacht, ein Link zur Lizenz beigefügt und angegeben werden, ob Änderungen vorgenommen wurden. Weitere Details zur Lizenz entnehmen Sie bitte der Lizenzinformation auf <https://creativecommons.org/licenses/by-nc-sa/4.0/deed.de>

In den Arbeitspapieren werden die Ergebnisse der Forschungsförderung der Otto Brenner Stiftung dokumentiert und der Öffentlichkeit zugänglich gemacht. Für die Inhalte sind die Autorinnen und Autoren verantwortlich.

Download und weitere Informationen:
www.otto-brenner-stiftung.de

Vorwort

Donald Trump verkündet Milliarden-Investitionen in die KI-Infrastruktur, Elon Musk will die Planung des Bundeshaushalts der USA künftig einer KI überlassen, das chinesische Start-up DeepSeek bringt mit seinem neuen KI-Modell die Aktienkurse des Chipherstellers Nvidia kurzzeitig zum Absturz und die großen Social-Media-Plattformen machen KI zum festen Bestandteil ihres Angebots, während sie zugleich ankündigen, den Wahrheitsgehalt von Inhalten zukünftig nicht mehr überprüfen zu lassen – seit der Veröffentlichung von Chat-GPT Ende 2022 reißen die Meldungen zum Thema Künstliche Intelligenz nicht ab. Die Technologie ist zum zentralen Instrument im Wettrüsten von Investoren, Tech-Unternehmen, Start-ups und Staaten geworden. Zugleich beflügelt der KI-Hype utopische Vorstellungen einer Zukunft, in der die Menschheit sich von nicht notwendiger Arbeit, Hunger und Krankheit befreit sowie dystopische Szenarien einer von Maschinen kontrollierten Welt permanenter Kontrolle und Überwachung.

Heiß diskutiert wird dabei auch, wie der rasante Fortschritt generativer KI-Modelle die Kampagnenarbeit politischer Akteur*innen sowie die Meinungsbildungsprozesse und Wahlentscheidungen der Bürger*innen verändert. Einigkeit besteht darin, dass die neue Technologie ein Gamechanger für den politischen Wettkampf der Parteien ist – strittiger ist, in welcher Weise dies geschieht. Während einige Stimmen den Einfluss der neuen Technologie für begrenzt halten, ist besonders in den Medien eine Neigung hin zu Gefahren-Szenarien zu beobachten, die das Missbrauchspotenzial der neuen KI-Modelle betonen und vor Manipulation, Desinformation und einer schwindenden Glaubwürdigkeit politischer Inhalte warnen.

Was aber denkt die deutsche Bevölkerung über den Einsatz von (generativer) KI in politischen Kampagnen? Wie werden KI-generierte Bilder erkannt und wahrgenommen? Unterscheiden sich KI-Bilder in ihrer Wirkung von „echten“ Fotografien? Diesen Fragen widmen sich die Kommunikationswissenschaftler*innen Simon Kruschinski, Pablo Jost, Hannah Fecher und Tobias Scherer von der Johannes Gutenberg-Universität Mainz. Sie haben rund 2.000 Menschen zu ihrem Wissen über KI befragt und um Einschätzungen zur Nutzung von KI in der Politik gebeten. Zudem haben sie in einem Online-Experiment die Wahrnehmung und Wirkung KI-generierter Kampagnenbilder untersucht.

Die Ergebnisse der repräsentativen Befragung zeigen, dass die skizzierten Gefahren-Szenarien Resonanz in der Bevölkerung finden: Gleichwohl die Befragten bisher

wenig Berührungspunkte mit künstlich erzeugten Inhalten in politischen Kampagnen hatten, stehen sie dem Einsatz von KI in der Politik mehrheitlich besorgt gegenüber. Diese Vorbehalte korrespondieren mit einem eindeutigen Wunsch nach verbindlichen Regeln zur Überprüfung und Transparenz von KI-generierten Inhalten.

Das Online-Experiment belegt Schwierigkeiten aufseiten der Teilnehmenden, zwischen KI-generierten und echten Kampagnenbildern zu unterscheiden. Während eine KI-Kennzeichnung bei der Erkennung hilft, führt diese auch dazu, dass nicht künstlich erzeugte Inhalte als KI-generiert eingestuft werden. Eine weitere zentrale Erkenntnis: In ihrer emotionalen Wirkung unterscheiden sich KI-Bilder kaum von realen Fotografien, beide Bildformen wirken über die dargestellte Stimmung sowie das Bildthema. Die Autor*innen schlussfolgern hieraus, dass (generative) KI die Wirkungsdynamiken in der politischen Willensbildung nicht grundlegend verändert. Ein Gamechanger ist die Technologie vielmehr, weil sie die Effizienz politischer Kampagnenarbeit steigert und eine gezielte Ansprache von Wählergruppen ermöglicht.

Mit Blick auf den aktuellen Bundestagswahlkampf lässt sich allerdings festhalten, dass zwar alle Parteien sporadisch KI-Bilder und Videos nutzen, die Kampagneninhalte mehrheitlich aber auf „klassische“ Weise produziert werden. Eine Ausnahme bildet die AfD, die KI deutlich häufiger einsetzt. Das „Aufregerthema KI“, welches mediale Berichterstattung garantiert, ermöglicht der Partei mit einfachen Mitteln viel Aufmerksamkeit zu erzeugen – bereits einige Male wurde über Wahlkampfaktionen der AfD nur deshalb berichtet, weil sie ihre völkischen Gewaltfantasien mittels KI-generierter Videos verbreitete. Zugleich versucht die Partei das allgemeine Misstrauen gegenüber politischen Inhalten zu verstärken. Je größer die Angst vor KI-generierten Desinformationen in der Wählerschaft ist, desto leichter lassen sich auch Zweifel an eigentlich wahren Aussagen und Bildern der politischen Gegner verbreiten. Ein Ziel, das die AfD mit Donald Trump teilt, der in seinem Schulterschluss mit großen Teilen des Silicon Valley daraufhin wirkt, den Einsatz generativer KI zu deregulieren und der politischen Kontrolle zu entziehen.

Ist es also ratsam, die Auswirkungen von KI auf die politische Meinungsbildung nicht zu überzeichnen, muss das Missbrauchspotenzial dieser Technik durch den neuen faschistischen Machtblock ernst genommen werden. Künstliche Intelligenz ist vorrangig kein technisches, sondern ein gesellschaftspolitisches Thema. Es braucht die kritische Auseinandersetzung, wie KI-Modelle entwickelt, angeboten, eingesetzt und politisch kontrolliert werden können. Stiftung und Autor*innen hoffen, zu dieser Debatte beizutragen.

Robin Koss

Referat Wissenschaftsförderung

Frankfurt am Main, im Februar 2025

Inhalt

1	Einleitung.....	5
1.1	Relevanz und Problemstellung.....	5
1.2	Zielsetzung und Aufbau der Studie	6
2	Setting the Scene: Generative KI in politischen Kampagnen	9
2.1	Generative KI: Was ist das, wie funktioniert sie und wie ist sie nutzbar?.....	9
2.2	Wie wird generative KI in politischen Kampagnen eingesetzt? Beispiele aus der Praxis.....	16
2.3	Wie verändert generative KI die politische Kommunikation? Chancen und Risiken	24
2.4	Regulierungsbedarf für politische KI-Kampagnen? Aktuelle Vorschläge und Diskussionen	28
2.5	Zwischenfazit: Der kontroverse Beginn einer technologischen Revolution.....	32
3	Befragung zum Einsatz von generativer KI in politischen Kampagnen.....	34
3.1	Erkenntnisinteresse und Forschungsfragen	34
3.2	Methode und Aufbau der Befragung.....	40
3.3	Ergebnisse der Befragung.....	42
3.4	Zwischenfazit: Wenig Berührungspunkte mit KI-Kampagnen, aber weit verbreitete Skepsis und Sorgen.....	50
4	Online-Experiment zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern	54
4.1	Erkenntnisinteresse und Forschungsfragen	54
4.2	Methode und Design des Experiments	62
4.3	Ergebnisse des Experiments	67
4.4	Zwischenfazit: Geringe Erkennung von KI-Bildern und Wirkungsdominanz der Bildinhalte über den KI-Ursprung	73
5	Fazit und Ausblick	77

6	Vorschläge für einen verantwortungsbewussten Umgang mit KI in politischen Kampagnen	83
6.1	Bürger*innen und zivilgesellschaftliche Akteur*innen: Förderung der Medienkompetenz im Umgang mit KI in politischen Kampagnen	83
6.2	Politische Akteur*innen: KI verantwortungsvoll in Kampagnen einsetzen	86
6.3	Journalist*innen: Über KI-Einsatz mit Wissen und Evidenz berichten	88
6.4	Regulierungsakteur*innen: Evidenzbasierte Richtlinien für den Einsatz von KI in politischen Kampagnen schaffen und durchsetzen	89
	Literatur	93
	Verzeichnis der Abbildungen.....	106
	Hinweise zu den Autor*innen	107

1 Einleitung

1.1 Relevanz und Problemstellung

In Indien wirbt ein verstorbener Staatsoberhaupt für die Unterstützung seines politischen Nachfolgers, in den USA überfällt Donald Trump eine Tankstelle mit einem Maschinengewehr und in Deutschland riegeln schwerbewaffnete Polizeibeamt*innen einen Weihnachtsmarkt mit Stacheldraht ab. All das ist „passiert“ – aber nichts davon ist real. Denn diese Ereignisse wurden für politische Kampagnen mithilfe sogenannter generativer künstlicher Intelligenz (KI) „erfunden“ und – komplettiert durch möglichst realistische Texte, Stimmen, Bilder oder Videos – an Wähler*innen über Social Media verbreitet.

Die zunehmende Nutzung von generativer KI in der politischen Kommunikation hat eine weltweite Debatte ausgelöst (Chowdhury, 2024; Dommett, 2023; Europäisches Parlament, 2024; Gillespie, 2024; Jungherr, 2023): Wie verändert ihr Einsatz die Art und Weise, wie politische Botschaften erstellt und verbreitet werden? Wie beeinflussen KI-generierte Botschaften die politische Meinungsbildung? Welche Möglichkeiten und welche Gefahren gehen vom Einsatz aus? Und wie kann sichergestellt werden, dass der Einsatz von KI in politischen Kampagnen transparent und verantwortungsvoll erfolgt?

Mit generativer KI wird eine Form der künstlichen Intelligenz bezeichnet, die audiovisuelle Inhalte wie Texte, Bilder, Videos, Musik oder Sprache produziert (OpenAI, 2016; Sanseviero et al., 2025). Dabei werden sogenannte generative KI-Modelle mit enormen Mengen an Daten trainiert und können auf dieser Grundlage schier unendliche Mengen an Inhalten in Sekundenbruchteilen erstellen (Gillespie, 2024; Goldstein et al., 2023). Diese sind oftmals nicht mehr von menschlicher Kreation zu unterscheiden und werden als sogenannte Deepfakes bezeichnet (Battista, 2024; Groh et al., 2024; Lu & Yuan, 2024).

In der Diskussion um die Möglichkeiten des Einsatzes und der Annahmen über die Wirkung generativer KI haben sich im gesellschaftlichen und wissenschaftlichen Diskurs drei wesentliche Bedenken herausgebildet: *Erstens* besteht die Gefahr einer Überflutung der öffentlichen Kommunikationsräume mit KI-generierten Inhalten, was insbesondere im Falle von Desinformationen zu einer Verschlechterung der Informationsqualität führen kann (Jungherr, 2023; Jungherr & Schroeder, 2023). *Zweitens* könnten KI-generierte, personalisierte Botschaften zur Polarisierung oder Fragmentierung der öffentlichen Meinungsbildung beitragen, wenn sie bereits bestehende Meinungen oder Überzeugungen verstärken

(Battista, 2024; Novelli & Sandri, 2024). *Drittens* besteht die Möglichkeit eines Vertrauensverlusts in die politische Kommunikation insgesamt, wenn Wähler*innen nicht mehr sicher sein können, ob eine Botschaft von einem Menschen oder einer Maschine erstellt wurde und ob das Abgebildete der Realität entspricht (Hameleers et al., 2024; McKay et al., 2024; Vaccari & Chadwick, 2020).

Diese Bedenken gegenüber generativer KI haben politische Entscheidungsträger*innen und Gesetzgeber*innen dazu veranlasst, Regeln und Gesetze vorzuschlagen, die den Einsatz einschränken (G'sell, 2024; Jungherr, 2024; Laude & Daum, 2024; McKay et al., 2024). So hat das Europäische Parlament mit dem „AI Act“ die weltweit erste umfassende Rahmengesetzgebung für den Einsatz von KI verabschiedet, die Verpflichtungen für Anbieter*innen und Nutzer*innen festlegt (Europäisches Parlament, 2024). Diese richten sich nach Risiken, die von KI-Systemen ausgehen. Über den Einsatz von KI-generierten Inhalten in politischen Kampagnen wird noch hitzig debattiert. Im Bundestagswahlkampf 2025 einigten sich alle Parteien außer der Alternative für Deutschland und dem Bündnis Sahra Wagenknecht in einem Fairnessabkommen darauf, KI-generierte Inhalte nur zu nutzen, wenn diese klar als solche gekennzeichnet sind und politischen Mitbewerber*innen keine Aussagen in den Mund legen (Spiegel.de, 2024). Einige deutsche Rechtsexpert*innen (u. a. Laude & Daum, 2024) gehen weiter und sehen den Einsatz als Verstoß gegen Artikel 21 Abs. 1 GG, dem nach Parteien zur Förderung des demokratischen Willensbildungs-

prozesses auf die Verbreitung richtiger Informationen achten müssen, was mit dem Einsatz generativer KI nicht gewährleistet sei.

Angesichts der Risiken aber auch der Potenziale der neuen Kommunikationstechnologien ist es von großer Bedeutung, wie Bürger*innen den Einsatz von KI-generierten Kampagnenbotschaften bewerten und wie sie KI-generierte Inhalte wahrnehmen. Denn auf ihre Einschätzung und Wahrnehmung von KI-Inhalten kommt es letztlich an, insbesondere wenn Social-Media-Plattformen wie Facebook, Instagram oder X künftig KI zum festen Bestandteil bei der Erstellung von Botschaften machen (Robertson, 2024; Sato, 2025; Wankhede, 2024) und den Wahrheitsgehalt von Inhalten nicht mehr durch unabhängige Organisationen oder Nachrichtenagenturen überprüfen lassen wollen (Kaplan, 2025). Dies gilt umso mehr in einer Zeit, in der Desinformationen zunehmend als Bedrohung für Demokratien angesehen werden (World Economic Forum, 2024). Zur skizzierten Debatte liefert die vorliegende Studie einen Beitrag, indem für den deutschen Kontext erste empirische Erkenntnisse zur Bewertung des Einsatzes generativer KI in politischen Kampagnen sowie zu deren Erkennung, Wahrnehmung und Wirkung auf Bürger*innen vorgestellt werden.

1.2 Zielsetzung und Aufbau der Studie

Unsere Studie knüpft an die aktuellen Debatten rund um den Einsatz und die Wirkung von generativer KI in politischen Kampagnen an und gliedert sich in die folgenden drei Teile:

Einsatz von generativer KI in politischen Kampagnen

Im ersten Teil geben wir einen Überblick über den Einsatz von generativer KI in der politischen Kampagnenkommunikation. Insbesondere die generativen KI-Tools *ChatGPT* und *Midjourney* werden zunehmend in politischen Kampagnen eingesetzt, um menschenähnliche Texte beziehungsweise fotorealistische Bilder zu erstellen (Bieber et al., 2024; Łabuz & Nehring, 2024; McKay et al., 2024; Schueler et al., 2024). Beide KI-Tools eröffnen für die politische Kampagnenkommunikation viele Möglichkeiten: Sie können die Kommunikation in vielen Bereichen effizienter, schneller und inklusiver machen. So können sie ohne großen Geld-, Zeit- oder Personalaufwand bei der strategischen Planung von Botschaften unterstützen, für die Korrektur, Übersetzung und Vereinfachung von Texten genutzt werden oder die Erstellung von maßgeschneiderten Wahlkampfbotschaften mit realistischen Bildmotiven erleichtern (Bai et al., 2023; Dommett, 2023; Novelli & Sandri, 2024; Schnöller, 2024). Andererseits hat der Einsatz von generativer KI in der politischen Kampagnenkommunikation auch das Potenzial zum Missbrauch: Kampagnen können täuschend echte kompromittierende Bilder, Videos oder Audioclips politischer Gegner*innen erstellen oder maßgeschneiderte Falschinformationen in Kommentaren, auf Internetseiten und Social Media an bestimmte Wähler*innengruppen verbreiten (Goldstein et al., 2023; Łabuz & Nehring, 2024; Martin, 2024; OpenAI, 2024b). Wir beschreiben die Einsatzmöglichkeiten von KI anhand aktueller Beispiele, zeichnen die Debatte

um Chancen sowie Risiken von KI in politischen Kampagnen nach und diskutieren aktuelle Regulierungsmaßnahmen.

In den anschließenden zwei Teilen der Studie stellen wir die empirischen Ergebnisse unserer Untersuchung zu den folgenden übergeordneten Forschungsfragen vor:

*Wie wird der Einsatz von generativer KI in politischen Kampagnen von Bürger*innen bewertet?*

Der zweite Teil liefert ein umfassendes Bild von den Einstellungen der deutschen Bevölkerung zum Einsatz von generativer KI in politischen Kampagnen. Dafür haben wir eine repräsentative Online-Befragung von knapp 2.000 Teilnehmenden durchgeführt und Menschen zu ihrem KI-Wissen und ihren Einschätzungen zur KI-Nutzung in der Politik befragt. Die Auswertung der Antworten liefert Ergebnisse dazu, wie in der deutschen Bevölkerung die Potenziale, Risiken und bisherigen Regulierungsansätze generativer KI in der politischen Kampagnenkommunikation wahrgenommen, eingeschätzt und bewertet werden und schließen damit eine bis dato bestehende Forschungslücke: Denn bislang gibt es kaum Wissen darüber, wie der Einsatz von generativer KI in der politischen Kommunikation von Bürger*innen bewertet wird.

Wie werden KI-generierte Kampagnenbotschaften wahrgenommen und welche Wirkungen haben sie auf Einstellungen zum Einsatz von KI?

Im dritten Teil der Studie untersuchen wir, inwieweit KI generierte Botschaften in politischen

Kampagnen von Bürger*innen erkannt und wahrgenommen werden, aber auch wie sie ihre Einstellungen hinsichtlich des politischen Einsatzes von generativer KI beeinflussen. Dafür haben wir ein Online-Experiment mit ebenfalls rund 2.000 neuen Teilnehmenden durchgeführt. Die Teilnehmenden wurden dabei zufällig verschiedenen Versuchsbedingungen zugeteilt, in denen sie entweder KI-generierte oder echte Kampagnenbilder sahen. Um mögliche Einflussfaktoren zu identifizieren, wurden systematisch das Thema, die Valenz der Texte und Bilder (positiv/negativ) und die Absender*in der Anzeigen sowie das Vorhandensein von Aufklärungstexten und KI-Kennzeichnungen variiert.

Die Ergebnisse des Experiments zeigen, inwiefern Bürger*innen KI-generierte Kampagnenbotschaften als solche erkennen, welche Emotionen sie beim Betrachten empfinden und wie sie die Anzeigen bewerten. Zudem erörtern wir, welchen Einfluss positive oder negative Botschaftsinhalte sowie KI-Kennzeichnungen und Aufklärungstexte auf die Wahrnehmung haben. Schließlich präsentieren wir Ergebnisse, wie sich die Konfrontation mit KI-generierten Kampagnenanzeigen auf die Einstellungen gegenüber KI in der Politik und den Parteien auswirkt, die diese Technologie nutzen. Mit diesen Ergebnissen erhalten wir erstmalig Antworten auf bisher offene Fragen zur Wahrnehmung und Wirkung von KI-generierten politischen Botschaften.

Abschließend diskutieren wir die Ergebnisse unserer Studie vor dem Hintergrund der aktuellen Debatten um die Risiken und Potenziale des Einsatzes von KI in der Politik. Hierüber leiten wir wichtige Implikationen für verschiedene Stakeholder*innen ab:

- *Bürger*innen* können durch die Studie ihr Wissen zum Einsatz und zur Wirkung von KI vertiefen und dieses Wissen im alltäglichen Umgang mit KI-generierten Inhalten nutzen.
- *Zivilgesellschaftlichen Akteur*innen* liefert die Studie Impulse für die Förderung von Medienkompetenz und digitaler Bildung im Kontext von KI und Politik.
- *Politischen Akteur*innen* bietet die Studie Anregungen für einen verantwortungsbewussten Einsatz von generativer KI in ihren Kampagnen.
- *Journalist*innen* finden Anhaltspunkte für eine evidenzbasierte Berichterstattung über den KI-Einsatz in der Politik ohne Alarmismus.
- *Regulierungsakteur*innen* erhalten eine empirische Grundlage für die Entwicklung von Richtlinien zum verantwortungsvollen Einsatz von KI in der politischen Kommunikation.

Unsere Studie soll zur breiteren gesellschaftlichen Debatte über die Rolle von KI in Demokratien beitragen. Sie soll dabei helfen, ein Bewusstsein für die Chancen und Risiken dieser Technologie zu schaffen und als Ausgangspunkt eines informierten öffentlichen Diskurses dienen.

2 Setting the Scene: Generative KI in politischen Kampagnen

Die Auseinandersetzung mit den Chancen und Risiken von generativer KI in politischen Kampagnen ist von zentraler Bedeutung für demokratische Wahlen und Diskurse. Einerseits bietet diese Technologie das Potenzial für eine effizientere und zielgerichtete politische Kommunikation, die Wähler*innen umfassender informieren und einbinden kann (Jungherr & Schroeder, 2023; Novelli & Sandri, 2024). Andererseits bergen KI-generierte Inhalte Risiken für die Authentizität und Transparenz des politischen Diskurses sowie die Gefahr einer Informationsüberflutung und möglichen Manipulation (Battista, 2024; Novelli & Sandri, 2024). Darüber hinaus ist eine Auseinandersetzung mit den Möglichkeiten und Herausforderungen von generativer KI in politischen Kampagnen notwendig, um angemessene rechtliche und ethische Rahmenbedingungen zu schaffen (G'sell, 2024; Jungherr, 2024) sowie einen breiten gesellschaftlichen Diskurs über die Rolle von Technologie in der Demokratie anzuregen (Dommett, 2023; Europäisches Parlament, 2024).

In diesem Teil der Studie wird ein grundlegendes Verständnis zum Einsatz von generativer KI in der politischen Kommunikation entwickelt. Zunächst geben wir einen Überblick über generative KI-Systeme und deren Funktionsweise. Erläutert werden insbesondere Large Language Models (LLMs) zur Texterstellung und Diffusions-Model-

le zur Bildgenerierung. Auch das Konzept des Promptings zur Erstellung von KI-generierten Inhalten wird eingeführt (2.1). Im zweiten Abschnitt stellen wir Beispiele dafür vor, wie generative KI schon jetzt in der politischen Kampagnenarbeit genutzt wird (2.2). Anschließend stellen wir dar, wie in Politik, Gesellschaft und Forschung derzeit die Chancen und Gefahren generativer KI diskutiert werden (2.3). Auch bisherige Vorschläge zur Regulierung des politischen KI-Einsatzes werden vorgestellt (2.4). Das Kapitel endet mit einem Zwischenfazit, in dem wir die zentralen Erkenntnisse zum Einsatz von generativer KI in politischen Kampagnen und deren Relevanz für unsere empirischen Untersuchungen diskutieren (2.5).

2.1 Generative KI: Was ist das, wie funktioniert sie und wie ist sie nutzbar?

Künstliche Intelligenz (KI) umfasst traditionell Computersysteme, die Aufgaben erledigen, für die normalerweise menschliche Intelligenz erforderlich ist, wie zum Beispiel Denken, Lernen und Entscheiden (Liu et al., 2022; Sanseviero et al., 2025). KI-Systeme werden mittels Algorithmen entwickelt, die darauf abzielen, menschliche kognitive Funktionen zu imitieren und Aufgaben ebenso gut oder besser als Menschen zu erledigen. Während KI einfache Aufga-

ben wie das Schreiben und Bearbeiten von Texten übernehmen kann, betonen Expert*innen, dass menschliche Überwachung und Eingriffe weiterhin notwendig sind, insbesondere bei komplexeren Aufgaben (Leaver & Srdarov, 2023; Wolfram, 2023).

Generative KI ist eine spezialisierte Form der künstlichen Intelligenz, die darauf abzielt, neue Inhalte wie Texte, Bilder, Videos und sogar Musik oder Sprache zu erstellen (Gruber & Votta, 2025; OpenAI, 2016; Sanseviero et al., 2025). Diese KI-Modelle werden mit großen Mengen an Trainingsdaten trainiert, um hierüber zu lernen, Muster, Kontexte und Zusammenhänge in diesen Daten zu erkennen. Basierend auf diesen Erkenntnissen können sie neue Inhalte generieren, die den gelernten Mustern entsprechen oder die Eigenschaften und Konzepte der Trainingsdaten nachahmen. Um Inhalte zu erstellen, können Menschen mit sogenannten Prompts, also Texteingaben oder -befehlen durch Nutzende, Anfragen an die KI-Modelle stellen (Phoenix & Taylor, 2024). Dies erfolgt in einfacher Sprache und in Dialogform, sodass die KI-Modelle ohne große Vorkenntnisse genutzt werden können. Beispiele für Entwicklungsunternehmen von generativen KI-Systemen sind unter anderem OpenAI, Meta, Google, DeepMind, Anthropic und Midjourney (OpenAI, 2016; Sætra, 2023; Wankhede, 2024).

Um Texte oder Bilder zu verarbeiten, greifen aktuelle generative KI-Systeme in der Regel auf zwei unterschiedliche Modelle zurück, die folgend kurz vorgestellt werden.

2.1.1 Large Language Modelle und Diffusionsmodelle

Large Language Modelle (LLMs; siehe auch Infobox 1) konzentrieren sich auf die Verarbeitung natürlicher Sprache. Sie sind in der Lage große Mengen an Text zu analysieren und zu lernen, wie Wörter und Sätze miteinander in Beziehung stehen (Gruber & Votta, 2025; Tunstall et al., 2023; Vaswani et al., 2023). Die „Wissensbasis“ dieser Modelle wurde mithilfe mehrerer Milliarden Dokumente, wie beispielsweise Internettexten, Gesetztestexten, wissenschaftlichen Publikationen oder Social-Media-Kommunikation trainiert. Auf dieser Basis können diese Modelle Spracheingaben verstehen sowie logische und kohärente Texte erzeugen. Dies hilft ihnen, Aufgaben wie Übersetzungen, das Beantworten von Fragen oder das Schreiben von Texten zu bewältigen (Lee, 2023).

Diffusionsmodelle konzentrieren sich auf die Erzeugung realistischer Bilder und Videos. Diese Modelle sind Algorithmen, die an großen Mengen von Bilddaten trainiert werden und lernen, anhand ihrer Struktur Muster und Details in Bildern zu erkennen (Sanseviero et al., 2025). Der Trainingsprozess dieser Modelle nutzt Milliarden von Bildern, die aus verschiedenen Quellen wie Fotografien, Kunstwerken und digitalen Grafiken stammen. Durch sie lernen Diffusionsmodelle, welche Muster und Strukturen bestimmte Objekte in Bildern besitzen und können diese entsprechend rekonstruieren. Besonders bei der Entwicklung solcher „Text-To-Image“-Modelle (TTI) wurden innerhalb der letzten Jahre große technische Fortschritte erzielt, sodass es ihnen aus Texteingaben möglich ist, Thema, Stil, Kon-

text oder Stimmungen zu verstehen und in Bilder zu übersetzen (Phoenix & Taylor, 2024). Ein Beispiel ist das KI-Tool Midjourney, das textuelle Eingaben auf Basis von LLMs versteht und darauf basierend Bilder mithilfe von Diffusionsmodellen erstellt (vergleichbare TTI-Tools sind Grok von X oder Flux von Black Forest Labs).

Eine Vielzahl an KI-Tools wurde bereits für die breite Öffentlichkeit zugänglich gemacht. ChatGPT (OpenAI), CoPilot (Microsoft) oder Gemini (Google), deren Fokus ursprünglich rein auf der Erstellung von Texten lag, können durch die Integration von TTI-Modellen inzwischen ebenfalls verwendet werden, um Bilder zu generie-



Infobox 1:

Large Language Model

Ein Large Language Model (LLM) ist ein Computerprogramm, das entwickelt wurde, um natürliche Sprache zu verstehen und zu erzeugen (Lee, 2023). Es handelt sich um Wahrscheinlichkeitsmodelle, die statistische Wort- und Satzfolge-Beziehungen aus einer Vielzahl von Textdokumenten gelernt haben. LLMs basieren auf Transformer-Architekturen (Vaswani et al., 2023), die sich im Wesentlichen in zwei Hauptkomponenten unterteilen lassen: den Encoder und den Decoder.

Der Encoder nimmt den Eingabetext und zerlegt ihn in kleine Teile, die Tokens genannt werden. Diese Tokens können einzelne Wörter oder sogar Teile von Wörtern sein. Zum Beispiel kann der Satz „Katze und Hund sitzen auf der Bank“ in die Tokens [„Katze“, „und“, „Hund“, „sitzen“, „auf“, „der“, „Bank“, „.“] zerlegt werden. Diese Tokens sind leichter und sparsamer für das Modell zu verarbeiten als ganze Sätze. Danach werden den Tokens Zahlenwerte zugeteilt, die sich an den umliegenden Tokens in einem Dokument orientieren und auf Basis der Trainingsdaten berechnet werden (beispielsweise mit welchen anderen Worten wird „Katze“ häufig verwendet). Diese sogenannten Embeddings geben den Tokens eine Bedeutung und liefern Informationen darüber, wie Buchstaben oder Wörter zueinander in Beziehung stehen. Dadurch kann einerseits anhand der Ähnlichkeit der Zahlenwerte die Bedeutung von Worten erkannt werden („Hund“ und „Katze“). Andererseits kann der Kontext von Wörtern besser verstanden werden („Bank“ im Zusammenhang mit „Katze“ als Sitzmöglichkeit und nicht als Finanzinstitut). Der Decoder nutzt dann diese Tokens und ihre Embeddings, um den Text und die Beziehungen zueinander zu verstehen und neuen Text zu generieren. Er erstellt Sätze, die auf den ursprünglichen Tokens und ihren Bedeutungen basieren. Dadurch können LLMs Texte erzeugen, die logisch und kohärent klingen.

ren (Sastry et al., 2024). Die Möglichkeiten von TTI-Modellen werden zunehmend auch in Bildbearbeitungsprogrammen wie Photoshop oder Canva integriert, was den Zugang zu und die Bedienung von den KI-Modellen erleichtert. Auch eine Erstellung von Musik, Stimmen oder Videos ist inzwischen mittels generativer KI wie Runway, Fliki, Sora, GoogleVevo oder KLING AI möglich (Mollick, 2024; Sætra, 2023).

Die rasante technische Entwicklung der letzten Jahre bedingt, dass die künstlich generierten Ausgaben inzwischen oftmals nicht mehr von menschlich erzeugten Inhalten unterschieden

werden können (siehe Frank et al., 2023; Groh et al., 2024; Köbis et al., 2021). Bei solchen sogenannten „Deepfakes“ bestehen allerdings Unterschiede in der Qualität der jeweiligen Ergebnisse, je nach gewähltem Tool oder der genutzten Technik (siehe Infobox 2). TTI-Modelle erlauben detaillierte Gestaltungsmöglichkeiten hinsichtlich einzelner Elemente wie abgebildeter Personen, Kamerawinkel oder Beleuchtung, sodass sie vermehrt für die Erstellung von Deepfakes genutzt werden (Knochel, 2023; Wankhede, 2024). Damit die KI-Bilder aber realistisch aussehen, müssen die richtigen Strategien bei den Texteingaben genutzt werden.

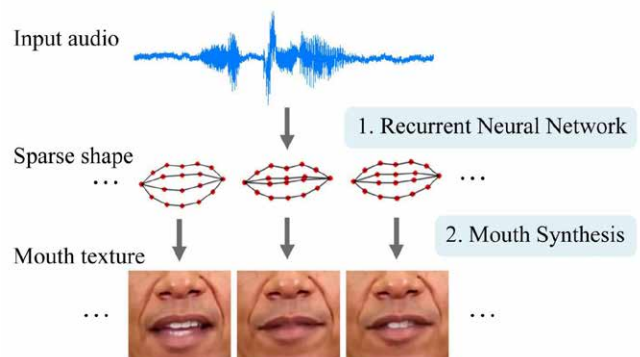


Infobox 2:

Was sind Deepfakes?

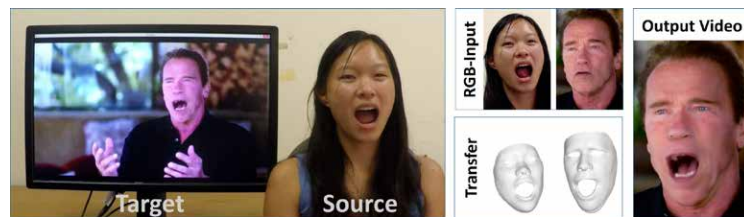
Deepfakes sind realistisch wirkende Bilder, Videos oder Audiodateien, die mithilfe künstlicher Intelligenz erstellt oder verändert wurden. Sie ermöglichen es, Objekte und Personen sowie ihre Gesichter, Stimmen und Bewegungen täuschend echt zu simulieren. Es gibt unterschiedliche Deepfake-Techniken, die mithilfe verschiedener Tools auch in Echtzeit und in Kombination eingesetzt werden können:

Lip Syncing: Diese Methode verändert die Lippenbewegungen einer Person in einem Video, sodass sie zu einem anderen Audiosignal passen (Suwajanakorn et al., 2017). Dabei werden die Lippen aus einem Originalbild und der Ton aus einer Audioaufnahme oder von einer Texteingabe aufeinander abgestimmt. Dadurch sieht es so aus, als ob die Person in einem Video tatsächlich den neuen Sprechertext verwendet.



Face Swapping: Face Swapping ist eine Technik, bei der das Gesicht einer Person in einem Video durch das Gesicht einer anderen Person ersetzt wird. Dies wird häufig in Social-Media-Apps verwendet, um lustige oder künstlerische Effekte zu erzielen, hat aber auch Anwendung in Deepfakes gefunden, um überzeugende Fälschungen zu erstellen.

Face Reenactment: Bei Face Reenactment wird die Mimik einer Person auf eine andere übertragen (Thies et al., 2020). Ein bekanntes



Beispiel ist die Technik „Face2Face“, bei der die Gesichtsausdrücke eines Live-Actors auf ein Zielvideo projiziert werden, um die Illusion zu erzeugen, dass die abgebildete Person tatsächlich die Bewegungen des Actors ausführt.

Motion Transfer: Motion Transfer überträgt die Bewegungen einer Person auf eine andere, einschließlich Gestik und Körperhaltung. Diese Technik ermöglicht es, die Bewegungen eines Tänzers oder Sportlers auf eine andere Person zu projizieren, wodurch realistische Fälschungen von Aktivitäten entstehen.

Prompt Engineering: Prompt Engineering ist eine Methode, bei der eindeutige Befehls-eingaben für KI-Tools erarbeitet werden, die es generativen KI-Modellen erleichtern, das von Nutzer*innen gewünschte Ergebnis zu erzeugen. Dadurch lassen sich unter anderem Inhalt, Stil, Stimmung und Natürlichkeit von Bildern steigern, sodass sie wie reale Fotografien aussehen.

Allgemeines Prompting: „Angela Merkel auf einer Couch“



Prompt-Engineering: „Aufnahme von Angela Merkel auf einer roten Couch im Sonnenlicht vor einer Wand mit einem Künstlerbild, fotorealistischer Stil“



Prompt-Engineering: „Aufnahme von Angela Merkel auf einer schwarzen Couch im Lampenlicht in einem alten Wohnzimmer, fotorealistischer Stil“



2.1.2 Prompting und Prompt Engineering

Die Eingaben von Befehlen durch Nutzer*innen in KI-Tools werden als „Prompts“ bezeichnet (Burkhardt & Rieder, 2024; Phoenix & Taylor, 2024). Diese können sowohl wenige als auch eine Vielzahl an Sätzen umfassen und geben die Ausgabe vor, die Nutzer*innen als Ergebnis von der generativen KI erwarten. Um exakt die gewünschte Ausgabe zu erhalten, bedarf es dabei sorgfältiger und präziser Eingaben, die auf spezifischen Logiken der KI-Modelle oder -Systeme basieren. Damit die generative KI Befehle optimal „verstehen“ kann, sollten Prompts so formuliert sein, dass sie präzise sind, einen Kontext erwähnen, Beispiele beinhalten und der KI

eine bestimmte Rolle zuweisen, beispielsweise Marketingexpert*in, Wissenschaftler*in oder Politiker*in (Oppenlaender, 2023; Phoenix & Taylor, 2024).

Das Erarbeiten spezifischer Eingaben wird auch „Prompt-Engineering“ genannt (siehe Infobox 2) und kann dabei helfen, gewünschte und qualitativ hochwertige Ergebnisse von einer generativen KI zu erhalten (Burkhardt & Rieder, 2024; Phoenix & Taylor, 2024). Dabei kann generative KI aus den Eingabedaten lernen und sich entsprechend an die gewünschte Ausgabe anpassen. Dies hilft, die Lücke zwischen einfachen Eingaben beziehungsweise Befehlen und sinnvollen Ausgaben

Abbildung 1

Struktur für das Prompt Engineering in *Midjourney*

Kameraeinstellungen & Betrachtungswinkel				
Einstellungsgröße	Kameraperspektive			
Vordergrund				
Beschreibendes Adjektiv/Verb	Menschen, Objekte, Ereignisse	Handlung	Beschreibung Umgebung	„Im Vordergrund“
Hintergrund				
Beschreibendes Adjektiv/Verb	Menschen, Objekte, Ereignisse	Handlung	Beschreibung Umgebung	„Im Hintergrund“
Übergeordnete Stimmung & Stilistik				
Beschreibende Adjektive	Bildstilistik	Atmosphäre	Komposition	
Fotorealistische Prompt Elemente				
Belichtung	Farbstil	Kameramarke	Kameraobjektiv	
Midjourney spezifische Parameter				
Bildverhältnisse	Charakterreferenz	Gewichtung	Referenz-Nummer	Modellversionen
Quelle: Eigene Darstellung.				

der generativen KI zu schließen. Bei einem effektiven Prompt-Engineering gelingt es Nutzer*innen also, technisches Wissen mit einer umfassenden Sprachkenntnis zu kombinieren, um so schnellstmöglich optimale Ergebnisse zu erzielen.

Im Falle des Prompt Engineering bei TTI-Modellen sollten Eigenschaften und Charakteristika von Bildern in den Prompt eingebunden werden,

um diesen zu spezifizieren (siehe Abbildung 1). Bei Midjourney¹ kann der Textprompt mit einer Anweisung für die Kameraeinstellung und den Betrachtungswinkel des Bildes beginnen (Knochel, 2023; Oppenlaender, 2023; Phoenix & Taylor, 2024). Anschließend folgt die Beschreibung, was im Vordergrund des Bildes zu sehen sein soll. Hier kann das zentrale Bildelement – zum Beispiel Personen, Objekte, Ereignisse – ge-

¹ Im Bericht konzentrieren wir uns auf Midjourney, weil es zum Zeitpunkt der Veröffentlichung die detailliertesten Optionen zum Erstellen von fotorealistischen KI-Bildern ermöglichte. Damit erfüllte Midjourney am besten die Anforderungen, um unser Online-Experiment mit möglichst realistischen KI-Bildern durchzuführen.

nannt und mit einer entsprechenden Emotion und/oder einer beschreibenden Handlung oder Umgebung präzisiert werden. Im nächsten Teil des Prompts kann der Hintergrund des zu generierenden Bildes mit derselben Systematik beschrieben werden. Im darauffolgenden Teil des Prompts wird die Stimmung und Stilistik eingegeben – beispielsweise Anweisungen zu Wetter, Tageszeit, Stimmung, Atmosphäre oder Komposition. Im vorletzten Teil können Elemente beschrieben werden, die für ein fotorealistisches Bild hilfreich sind – so kann die Belichtung, der Farbstil, die Kameramarke und das Kameraobjektiv definiert werden. Abschließend können auf Midjourney noch Bildverhältnisse, Referenzen für konsistente Charakter und die zu verwendende KI-Modellversion beschrieben werden.

2.2 Wie wird generative KI in politischen Kampagnen eingesetzt? Beispiele aus der Praxis

Generative KI hat längst Einzug in die politische Kampagnenkommunikation gehalten. Von Wahlplakaten über Social-Media-Inhalte bis hin zu Deepfake-Videos – die Technologie wird weltweit von politischen Akteur*innen für verschiedenste Kommunikationszwecke eingesetzt. Während einige Parteien und Organisationen transparent mit der KI-Nutzung umgehen, setzen andere die Technologie ohne Kennzeichnung oder sogar gezielt zur Desinformation ein. Im Folgenden geben wir einen Überblick über

konkrete Beispiele der KI-Nutzung in politischen Kampagnen – von Akteur*innen jenseits Europas über deutsche Parteien bis hin zu deutschen zivilgesellschaftlichen Akteur*innen. Dabei erheben wir keinen Anspruch auf Vollständigkeit noch auf Repräsentativität, sondern zeigen lediglich Beispiele für die vielfältigen Einsatzmöglichkeiten von KI in aktuellen politischen Kampagnen, um die damit verbundenen Chancen und Herausforderungen für demokratische Prozesse zu verdeutlichen.²

Politische Kampagnen außerhalb von Deutschland haben generative KI schnell nach der Veröffentlichung von ChatGPT und Midjourney in ihrer Kampagnenkommunikation genutzt.

In den USA setzen sowohl Akteur*innen der republikanischen als auch der demokratischen Partei generative KI ein: Die Republikaner veröffentlichten beispielsweise zu Beginn des US-Präsidentschaftswahlkampfes 2023 ein ausschließlich mit KI generiertes Video, das bedrohliche Szenarien für die USA zeigt, falls Joe Biden als Präsident wiedergewählt werde (Der Standard, 2023). Das Video enthielt einen Hinweis auf die KI-Nutzung (Der Standard, 2023). Floridas republikanischer Gouverneur Ron DeSantis nutzte generative KI für Angriffe auf Donald Trump während des republikanischen Vorwahlkampfes, beispielsweise eine künstlich erstellte Stimme oder KI-generierte Bilder von Trump (siehe Abbildung 2 links; Harper et al., 2023).

2 Eine sich aktualisierende Sammlung zum Einsatz von KI in politischen Kampagnen auf der Welt findet sich beim „The WIRED AI Election Project“ (Elliott, 2024).

Kampagnen der Demokratischen Partei setzten generative KI im US-Präsidentenwahlkampf vor allem ein, um Wähler*innengruppen gezielt anzusprechen. Dafür wurden beispielsweise KI-gestützte Chatbots auf Plattformen wie WhatsApp, Facebook Messenger, Discord oder Twitch genutzt, um mit jungen Latinx und Afroamerikaner*innen ins Gespräch zu kommen (Carrasquillo, 2024).

Darüber hinaus verbreiten Social-Media-Nutzer*innen und politische Aktivist*innen KI-generierte Stimmnahmen, Bilder und Videos von US-amerikanischen Politiker*innen. So gerieten KI-generierte Bilder von Kamala Harris als Kommunistin oder von Taylor-Swift-Fans als Trump-Unterstützer*innen in Umlauf, die auf Donald Trumps Social-Media-Accounts weiter-

verbreitet wurden (Robins-Early, 2024). Außerdem wurden KI-generierte Videos veröffentlicht, die Politiker*innen wie Joe Biden, Barack Obama, Donald Trump oder Kamala Harris beim Begehen von Straftaten zeigen (siehe Abbildung 2 rechts; Prosser, 2024).

Außerhalb der USA gibt es weitere Beispiele für den Gebrauch von KI in politischen Kampagnen. In *Indonesien* nutzte der Gewinner der Präsidentschaftswahl 2024, Prabowo Subianto, generative KI, um sich im Gewand einer niedlichen computeranimierten Figur zu zeigen (siehe Abbildung 3 links). Dadurch sollte Kritik in Bezug auf seine Rolle in der vergangenen Militärdiktatur des Landes geschwächt werden (Chowdhury, 2024). In *Indien* lässt sich ebenfalls ein zunehmender Gebrauch von generativer KI im Wahlkampf be-

Abbildung 2:

KI-generiertes Bild von Donald Trump und Dr. Anthony Fauci (links) und KI-generiertes Überwachungsvideo von Kamala Harris beim Begehen einer Straftat (rechts)



Quellen: Wodecki, 2023 / X/The Dor Brothers

obachten, vornehmlich zur Erstellung von Bildern und Videos der Kandidat*innen (Shah, 2024). Allerdings wird KI auch zur Simultanübersetzung von Reden des Premierministers Narendra Modi in bis zu 14 verschiedene indische Sprachen genutzt (Universum, 2024). Auch in *Pakistan* kam generative KI in politischen Kampagnen zum Einsatz. So erstellte das Wahlkampfteam des inhaftierten Kandidaten Imran Khan ein Video, in dem er sich mit einer KI-geklonten Stimme an die Bevölkerung wandte (Hegewisch, 2024; Universum, 2024).

Weitere Beispiele für den Einzug von KI in politische Kampagnen lassen sich in *Europa* finden. Beispielsweise warben mehrere *französische Parteien* bei der Parlamentswahl 2024 über Social Media mit KI-generierten Inhalten ohne diese zu kennzeichnen (Schueler et al., 2024). Insbesondere Rassemblement National, Reconquête und Les Patriotes machten damit Werbung

gegen Europa und Immigration (siehe Abbildung 3 rechts). Auch in der *Schweiz* nutzte die Partei „FDP.Die Liberalen“ bereits im Sommer 2023 KI zum Generieren politischer Kampagnen, etwa für ein Bild, das Klimaaktivist*innen bei einer Straßenblockade zeigte (Aregger, 2023). Im Wahlkampf zum Schweizer Nationalrat 2023 erstellte der sozialdemokratische Kandidat Islam Alijaj, der an Zerebralparese leidet, einen Video-Avatar, der auf Basis generativer KI seine Worte in eine synthetische Stimme umwandelt, die für alle verständlich ist (Fuchs, 2024).

Neben dem Einsatz von etablierten politischen Akteur*innen wird KI auch zunehmend in *verdeckten politischen Kampagnen* eingesetzt. So deckte OpenAI mehrere Kampagnen aus Russland, Israel, China und Iran auf, die ChatGPT nutzten, um Internetseiten von etablierten (inter)nationalen Medienanbietern nachzuahmen oder Kommentare auf Social Media zu erstellen (OpenAI, 2024b).

Abbildung 3:

KI-generierte Avatare von Prabowo Subianto und seines „Running Mate“ Gibran Rakabuming Raka (links) und KI-generierte Social-Media-Posts der Reconquête (rechts)



Quellen: Tan & Husada, 2024 / Schueler et al., 2024.

Die erstellten Inhalte wurden insbesondere in den USA und Europa verbreitet und zielten auf aktuelle politische Debatten zum Krieg in der Ukraine und zur Einwanderungs- oder Energiepolitik ab.

Und wie sieht die Nutzung von generativer KI in deutschen Kampagnen aus? In *Deutschland* wird generative KI größtenteils noch experimentell und nicht sehr strategisch von Parteien, ihren Kandidierenden oder anderen politischen Akteur*innen in Kampagnen eingesetzt, was die folgenden Beispiele verdeutlichen.³

Mit Blick auf die *SPD* lässt sich ein erster Gebrauch von generativer KI finden. Im Thüringer Landtagswahlkampf 2024 postete die Thüringer SPD einen CDU-Wahlkampfsport auf Social Media, in dem mit KI die Stimmen der Prota-

gonist*innen und des CDU-Spitzenkandidaten Mario Voigt verändert wurden (Görmann, 2024). Zudem verwendete die nordrhein-westfälische SPD-Landtagsfraktion während einer aktuellen Stunde des Landtages ein KI-generiertes Bild, das ein Mädchen in Anlehnung an Greta Thunberg zeigte (Auster, 2024; siehe Abbildung 4). Auch der baden-württembergische SPD-Landesverband nutzte KI, um Logos und Flyer zu erstellen (Meisoll, 2024). Darüber hinaus setzten auch einzelne SPD-Politiker*innen KI ein. Im Dezember 2022 wandte sich der forschungspolitische Sprecher der Brandenburger SPD-Landtagsfraktion, Erik Stohn, als erster Politiker öffentlichkeitswirksam mithilfe von ChatGPT an die Landesregierung und stellte so eine kleine Anfrage (Fuchs, 2024). Auch erstellte der bayerische SPD-Politiker Lars Mentrup Wahlplakate mithilfe

Abbildung 4:

KI-generiertes Bild der SPD-Landtagsfraktion NRW



Quelle: Auster, 2024.

³ Eine sich aktualisierende Sammlung von KI-generierten Kampagneninhalten, die von deutschen politischen Akteur*innen auf Social Media verbreitet werden, findet sich im „CampAlign Tracker“ (Kruschinski et al., 2025).

des KI-Modells DALL-E (Schubert, 2023; Warrlich, 2023) und der SPD-Bundestagsabgeordnete Bengt Bergt teilte ein mithilfe von KI manipuliertes Video von Friedrich Merz auf seiner Instagram-Seite (Mühlenkamp, 2024).

Bei der *CDU/CSU* gibt es ebenfalls Beispiele für die Nutzung von KI: So setzte die CDU Sachsen in ihrer Briefwahlkampagne bei der Landtagswahl 2024 ein KI-generiertes Bild von einem Mann beim Grillen ein, ohne dies zu kennzeichnen (Töpfer, 2024; Abbildung 5 links). Im Rahmen der Fußball-Europameisterschaft verwendete der nordrhein-westfälische Ministerpräsident Hendrik Wüst KI, um eine Videobotschaft in elf Fremdsprachen zu synchronisieren (t-online, 2024). Auch stellte ein CDU-Abgeordneter eine Anfrage an die baden-württembergische Landesregierung mittels KI (Meisoll, 2024). Die Schwes-

terpartei CSU nutzt KI bislang vor allem zur Erstellung von grafischen Elementen in der Social-Media-Kommunikation (Huesmann, 2024).

Als Beispiele für die KI-Nutzung durch die *FDP* kann die Erstellung von Logos und Flyern in Baden-Württemberg genannt werden (Meisoll, 2024). KI-generierte Bilder wurden aber auch auf den Social-Media-Kanälen der nordrhein-westfälischen FDP-Landtagsfraktion zur Themensetzung und für Kritik an der Landesregierung genutzt – und als KI gekennzeichnet (Auster, 2024, siehe Abbildung 5 rechts). Auch im Wahlkampf-Rap „Hessen auf die 1“ des hessischen Landtagsabgeordneten Yanki Pürsün wurde die Begleitstimme mit KI erstellt (Fuchs, 2024). Darüber hinaus setzte die FDP Hessen einen KI-Chatbot auf ihrer Webseite ein, der Fragen zum Wahlprogramm beantworten sollte (ebd.).

Abbildung 5:

KI-generierte Social-Media-Posts der CDU Sachsen (links) und der FDP-Landtagsfraktion NRW (rechts)



Quellen: Facebook/CDU Sachsen, FDP-Landtagsfraktion NRW.

Bei der *AfD* gebrauchen mittlerweile viele Parteigliederungen und -vertreter*innen KI strategisch für Kampagnen auf Social Media, Internetseiten und Plakaten, wobei größtenteils auf eine Kennzeichnung verzichtet wird. So veröffentlicht die Bundes-AfD auf ihren Social-Media-Kanälen und Internetseiten regelmäßig KI-generierte Bilder zum Thema Migration, Kriminalität oder innere Sicherheit. Auch in den einzelnen Bundesländern, wie in Sachsen, Brandenburg, Baden-Württemberg oder Rheinland-Pfalz, werden KI-generierte, nicht gekennzeichnete Inhalte auf den Social-Me-

dia-Accounts der Landesparteien oder -fraktionen genutzt. Darunter Bilder und Videos von stereotypen deutschen Familien und Männern, aggressiven Migrant*innen, Weihnachtsmärkten hinter Stacheldraht oder vom vermeintlich zerstörten Feldberg (siehe Abbildung 6; Rottach, 2024).

AfD-Kreis- und Ortsverbände nutzen ebenfalls KI für ihre Social-Media-Kommunikation. So stellte der Kreisverband Göppingen mittels generativer KI Personen, die in Social-Media-Posts mit Namen vorgestellt wurden und Gründe für

Abbildung 6:

KI-generierte Social-Media-Posts der AfD-Fraktion Baden-Württemberg (oben links), der AfD-Fraktion Rheinland-Pfalz (oben rechts) und Screenshots aus KI-generierten Video der AfD Brandenburg (unten)



Quellen: Facebook/AfD-Fraktion BaWü, AfD-Fraktion RLP, AfD Brandenburg.

ihren Parteieintritt nannten (Huesmann, 2024). Ebenso finden sich KI-generierte Bilder auf den Social-Media-Accounts von AfD-Politiker*innen. So postete Norbert Kleinwächter nicht-kennzeichnete KI-Bilder von blutverschmierten Mädchen, aggressiven Migrant*innen oder überzeichneten Regierungspolitiker*innen (Lauer, 2023). Darüber hinaus nutzte die AfD Sachsen KI für einen Chatbot, der Fragen zum Wahlprogramm beantworten sollte (Scholl, 2024). Als weiteres Beispiel ist der „Ampel-Adventskalender“ zu nennen, der unter anderem über YouTube- und TikTok-Kanäle der Partei veröffentlicht wurde (Laude & Daum, 2024). Dabei wurden täglich mit KI manipulierte Sprachnachrichten von Mitgliedern der Ampel-Bundesregierung um Olaf Scholz, Robert Habeck oder Annalena Baerbock veröffentlicht.

Abbildung 7:

KI-generierter Social-Media-Post
der Linken Sachsen



Quelle: Facebook/Die Linke Sachsen

Deutlich weniger Beispiele für einen KI-Einsatz lassen sich für *Bündnis 90/Die Grünen*, *Die Linke* und das *Bündnis Sahra Wagenknecht* finden. So wurde KI für „Kreativprozesse“ bei der Erstellung von Social-Media-Inhalten durch den Bayerischen Landesverband der *Grünen* verwendet (Fuchs, 2024) und Pressemitteilungen in Baden-Württemberg erstellt (Meisoll, 2024). Im baden-württembergischen Landtag haben die Grünen zudem eine Rede mittels KI schreiben lassen und vorgelesen (Meisoll, 2024). Dagegen nutzte *Die Linke* in Sachsen KI, um grafische Elemente für Social-Media-Posts über soziale Ungerechtigkeit oder mit Kritik an der Sächsischen Landesregierung zu verfassen – und kennzeichnete diese als solche (Fuchs, 2024, siehe Abbildung 7). Darüber hinaus half KI dem Linken Bayerischen Landesverband Bilder und Formulierungen auf Wahlplakaten und Flyern zu verbessern (Schubert, 2023).

Für das *Bündnis Sahra Wagenknecht* gab es zum Recherchezeitpunkt kein Einsatzbeispiel von generativer KI in der Parteikommunikation. Die Partei gab im Gespräch mit dem Redaktionsnetzwerk Deutschland lediglich an, sich noch nicht auf die Verwendung dieser Technologie im Wahlkampf festgelegt zu haben (Huesmann, 2024).

Generative KI findet darüber hinaus auch Verwendung bei anderen politischen Akteur*innen wie Behörden, NGOs oder Aktivist*innen. Beispielsweise veröffentlichte die Berliner Senatsverwaltung für Stadtentwicklung eine Kampagne, die KI-generierte Gesichter von Wohnungssuchenden zeigte (Roelcke, 2023). Dadurch sollte auf die angespannte Lage auf dem Berliner Wohnungsmarkt verwiesen und die Akzeptanz für

Neubauprojekte erhöht werden. Aus dem Kreis der NGOs veröffentlichte etwa der WWF mehrere KI-generierte Bilder zur Kampagne #WorldWithoutNature, um auf den Verlust der Artenvielfalt und Schädigungen der Umwelt durch den Menschen hinzuweisen (OMF, 2024). Weitere Verwendungen von generativer KI finden sich in Kampagnen von politischen Aktivist*innen. Ein Beispiel ist das Zentrum für Politische Schönheit, das nicht-gekennzeichnete KI-generierte Bilder von verhafteten AfD-Politiker*innen auf Internetseiten, Social Media und Berliner Installationen nutzte, um auf ein gefordertes Verbot der Partei

aufmerksam zu machen (Das Zentrum für Politische Schönheit, 2024; siehe Abbildung 8 oben). Die Aktivist*innen nutzten auch KI für ein Video von Olaf Scholz, in dem er die angebliche Einleitung eines Verbotsverfahrens gegen die AfD ankündigte. Auch die Initiative „AfDnee“ verwendete für ihre Kampagne im Internet und auf Plakaten KI-generierte Bilder von fiktiven ehemaligen AfD-Wähler*innen, um für Spenden zu werben und hypothetische Szenarien für den Fall einer Beteiligung der AfD an Landes- oder Bundesregierungen darzustellen (AfDnee, 2024, siehe Abbildung 8 unten). Auch vermeintlich private

Abbildung 8:

KI-generierte Bilder der Kampagne „AfD-Verbot“ vom Zentrum für Politische Schönheit (oben) sowie der Kampagne „AfDnee“ (unten)



Quellen: afd-verbot.de / afdnee.de

Accounts nutzen generative KI für politische Kampagnen oder Satire. Beispielsweise verbreitete sich ein mit generativer KI unterstütztes Video, in dem sich Tagesschau-Moderator*innen für ihre „dreisten Lügen“ entschuldigten (Reveland & Siggelkow, 2023). Andere KI-generierte Videos von deutschen Politiker*innen in Mittelalterszenarien oder Musikvideos wurden ebenfalls über Social Media verbreitet (Breschendorf, 2024).

2.3 Wie verändert generative KI die politische Kommunikation? Chancen und Risiken

Die im vorangegangenen Abschnitt diskutierten Beispiele zeigen, dass (generative) KI mittlerweile in vielen Kampagnen politischer Akteur*innen zum Einsatz kommt (Battista, 2024; Chowdhury, 2024; Dommert, 2023; Jungherr, 2023; Łabuz & Nehring, 2024). Diese Nutzung birgt sowohl vielversprechende Potenziale als auch erhebliche Risiken für demokratische Prozesse. Während die Technologie einerseits die Effizienz und Reichweite politischer Kommunikation steigern und den Zugang zur politischen Debatte erleichtern kann, bestehen andererseits Bedenken hinsichtlich Manipulation, Desinformation und eines möglichen Vertrauensverlusts in politische Akteur*innen. Im Folgenden analysieren wir diese Chancen und Risiken des KI-Einsatzes in der politischen Kampagnenkommunikation.

2.3.1 Chancen des Einsatzes von generativer KI

Weil generative KI über kostengünstige und einfach zu bedienende Tools für alle verfügbar ist, kann der Einsatz Kampagnen bei vielfältigen

Aufgaben effektiv und effizient unterstützen. So ermöglichen ChatGPT, CoPilot oder Claude beispielsweise die Übernahme von Recherchen, Zusammenfassungen oder das Erstellen, Übersetzen oder Korrekturlesen von Presstexten oder Social-Media-Posts (Gillespie, 2024; Hügelmann, 2024; Novelli & Sandri, 2024). Mit TTI-Modellen wie Midjourney können auch visuelle Kampagneninhalte für Wahlplakate, Social Media oder Webseiten erstellt werden, wodurch eine ansprechende Gestaltung erleichtert wird. Durch die technologischen Fortschritte sind die KI-generierten Inhalte oftmals nicht mehr von Texten und Bildern professioneller Kampagnenberatungen oder -agenturen zu unterscheiden (Frank et al., 2023; Groh et al., 2024; Knochel, 2023; Mollick, 2024). Zusätzlich sind inhaltliche und thematische Anpassungen nach parteipolitischer Leitlinie, Ideologie oder Zielgruppe mithilfe des Prompt Engineerings möglich (siehe Abschnitt 2.1.2; Phoenix & Taylor, 2024). So kann generative KI dazu verwendet werden, bestimmte Zielgruppen mit maßgeschneidereten Inhalten anzusprechen und Anzeigen entsprechend zu personalisieren und platzieren (Hackenburg & Margetts, 2023; LaChapelle & Tucker, 2023). Letztlich können die Inhalte in schier unendlichen Massen und größtenteils vollautomatisiert produziert werden, was zur Effizienzsteigerung und Ressourcenoptimierung beitragen kann (Hügelmann, 2024; Jungherr & Schroeder, 2023; Novelli & Sandri, 2024). Dies ermöglicht es auch ressourcenschwächeren Parteien oder Kandidat*innen, professionelle Kampagnen zu führen und somit ihre Reichweite und Wettbewerbsfähigkeit zu erhöhen. Im Umkehrschluss kann KI aber natürlich auch

ressourcenstarken Akteur*innen helfen, ihren strategischen Vorsprung gegenüber den „kleineren“ politischen Gegner*innen auszubauen (Dommett, 2023).

Für die demokratische Auseinandersetzung bietet der Einsatz von generativer KI das Potenzial, den politischen Diskurs zu bereichern und inklusiver zu gestalten. Die Technologie kann dazu beitragen, politische Themen verständlicher aufzubereiten und einer breiteren Öffentlichkeit zugänglich zu machen (Jungherr, 2023; Jungherr & Schroeder, 2023; Novelli & Sandri, 2024). Durch die Möglichkeit, Inhalte in verschiedenen Sprachen und Formaten zu erstellen, können unterrepräsentierte Bevölkerungsgruppen oder Nicht-Muttersprachler*innen besser erreicht und in den politischen Prozess eingebunden werden. Zusätzlich können personalisierte Botschaften für diese Bevölkerungsgruppen erstellt werden, die auf die Interessen und politischen Einstellungen von Wähler*innen zugeschnitten sind (Hackenburg & Margetts, 2023; LaChapelle & Tucker, 2023). Durch das gezielte Adressieren hätte dies das Potenzial, das politische Interesse und die politische Beteiligung zu fördern (Jungherr, 2023; LaChapelle & Tucker, 2023; Novelli & Sandri, 2024). Zudem kann die KI-gestützte Analyse großer Mengen an Kommunikation oder Anliegen von Bürger*innen helfen, responsiver auf die Bedürfnisse der Wähler*innen einzugehen (sog. Social Listening; Martin, 2024). Allerdings setzt dies einen verantwortungsvollen und transparenten Einsatz der Technologie voraus, um das Vertrauen in den demokratischen Prozess zu wahren und möglichen Missbrauch zu verhindern.

2.3.2 Risiken des Einsatzes von generativer KI

Es liegt auf der Hand, dass der Einsatz von generativer KI in politischen Kampagnen neben Chancen auch erhebliche Risiken und Herausforderungen birgt. Ein zentrales Problem ist der potenzielle Verlust von Arbeitsplätzen in der Kreativbranche, da viele Aufgaben zukünftig maschinell erledigt werden könnten (Huh et al., 2023; Short & Short, 2023). Zudem ergeben sich rechtliche Bedenken hinsichtlich möglicher Urheberrechtsverletzungen und das Risiko einer Verstärkung von Stereotypen und Vorurteilen durch KI-generierte Inhalte (Bianchi et al., 2023; Cao & Kosinski, 2024; Gillespie, 2024). Dies basiert darauf, dass KI-Systeme mit öffentlich zugänglichen – und auch urheberrechtlich geschützten – Daten trainiert werden und sie die darin vorhandenen Vorurteile und Verzerrungen lernen (Bianchi et al., 2023; Feng et al., 2023; Geva et al., 2019; Urman & Makhortykh, 2023; siehe auch Infobox 3 unten). Dies könnte dazu führen, dass Politiker*innen und Parteien unbeabsichtigt Vorurteile über Minderheiten verstärken (LaChapelle & Tucker, 2023). Um diesen Risiken zu begegnen, empfehlen Expert*innen, generative KI vorwiegend für Routineaufgaben einzusetzen und nicht für die Massenproduktion von Kampagnenmaterial (Dommett, 2023). Darüber hinaus ist ein transparenter Umgang mit KI-generierten Inhalten unerlässlich, um das Vertrauen der Öffentlichkeit in politische Akteur*innen zu wahren (Gordon, 2024). Andernfalls besteht das Risiko, dass Wähler*innen sich hinter das Licht geführt fühlen, was dem Ansehen von Parteien und Politiker*innen schaden und die Glaubwürdigkeit politischer Botschaften insgesamt verringern könnte (Gillespie, 2024; Novelli & Sandri, 2024; Sætra, 2023).

Ein häufig übersehener Aspekt des KI-Einsatzes in politischen Kampagnen ist der enorme Ressourcenverbrauch, der mit dem Training und der Nutzung generativer KI-Systeme einhergeht. Die Erstellung von KI-generierten Bildern, Videos und Texten benötigt erhebliche Rechenleistung und verbraucht entsprechend große Mengen an Energie und damit auch Geld. So zeigen Studien, dass das Training eines einzigen großen KI-Modells mehrere hunderttausend Kilogramm CO₂ produzieren kann – etwa so viel wie 125 Transatlantikflüge (Strubell et al., 2020; Rohde et al., 2021). Auch der Wasserverbrauch für die Kühlung der Rechenzentren ist beträchtlich: Die Generierung eines einzelnen KI-Bildes kann mehrere hundert Milliliter Wasser benötigen (Rohde et al., 2021). Dies spiegelt sich auch in den monetären Kosten wider. So war OpenAI im Jahr 2024 (noch) nicht profitabel und erwartete Verluste von etwa 5 Milliarden US-Dollar bei einem Umsatz von 3,7 Milliarden US-Dollar (Wiggers, 2025). Angesichts dieser Kosten müssen sich politische Akteur*innen kritisch fragen, ob der massenhafte Einsatz von KI-generierten Inhalten in ihren Kampagnen wirklich vertretbar ist – besonders wenn sie sich gleichzeitig für Klimaschutz und Nachhaltigkeit einsetzen. Dies unterstreicht die Notwendigkeit eines bewussten und selektiven Einsatzes dieser Technologie.

Auch für demokratische Prozesse stellt die Verwendung von generativer KI in der politischen Kommunikation eine Herausforderung dar. Eine zentrale Sorge ist die potenzielle Überflutung öffentlicher Kommunikationsräume mit KI-generierten Inhalten, was zu einer Verschlechterung

der Informationsqualität führen kann (Gillespie, 2024; Jungherr, 2023; Jungherr & Schroeder, 2023). Als besonders problematisch wird dabei angesehen, dass die KI-Inhalte oft schwer als Fälschungen zu erkennen sind (u. a. Frank et al., 2023; Groh et al., 2024; Köbis et al., 2021). Solche Deepfakes (siehe Infobox 2) bergen erhebliche Risiken, darunter die Verbreitung von Desinformation, Rufschädigung, und die Verletzung der Privatsphäre öffentlicher Personen und Politiker*innen (Battista, 2024; Bundesrat, 2024; Hameleers et al., 2024). Die Möglichkeit, schnell und gezielt große Mengen falscher oder irreführender Informationen zu erstellen, gefährdet die Integrität des politischen Diskurses (Jungherr & Schroeder, 2023; Vaccari & Chadwick, 2020). Zudem könnten KI-generierte, personalisierte Botschaften zur Polarisierung oder Radikalisierung von Diskursen beitragen, indem sie den Raum für einen ausgewogenen Meinungs austausch einengen und den Wahlkampf emotionaler gestalten (Dobber et al., 2021; Jungherr & Schroeder, 2023; Novelli & Sandri, 2024). Ein weiteres Problem ist der potenzielle Vertrauensverlust in die politische Kommunikation, wenn die Authentizität von Botschaften nicht mehr eindeutig feststellbar ist: Wenn Wähler*innen nicht mehr unterscheiden können, ob eine Botschaft von einem Menschen oder einer Maschine erstellt wurde, könnte dies zu einer generellen Skepsis gegenüber allen politischen Äußerungen führen (Hameleers et al., 2024; McKay et al., 2024; Vaccari & Chadwick, 2020). Dies kann von politischen Akteur*innen ausgenutzt werden, indem sie „echte“ Aussagen oder Bildmaterialien als Fälschungen abtun, um sich aus der Verantwortung für problematische Äußerungen



Infobox 3:

Algorithmischer Bias und Urheberrechtsverletzungen

Eine der mestdiskutierten Herausforderungen von künstlicher Intelligenz ist algorithmischer Bias. Maschinen werden häufig pauschal als frei von menschlichen Fehlern, wie Voreingenommenheit (auch Bias) wahrgenommen (Sundar, 2008). Dabei stellt diese Fehleinschätzung eine der zentralsten Herausforderungen von KI dar. Algorithmischer Bias entsteht einerseits durch die programmierenden Personen und menschliche Annotator*innen (Geva et al., 2019), andererseits durch Bias in den Trainingsdaten (Feng et al., 2023). So basieren Trainingsdatensätze für KI-Modelle auch auf großen Mengen an ungefilterten Internetquellen, die neben sozialem Bias auch strafbare Inhalte oder falsche Informationen enthalten (LaChapelle & Tucker, 2023).

Dadurch können KI-Modelle in der Gesellschaft bereits vorhandene Stereotypen, soziale Hierarchien und falsche Informationen in Trainingsdaten aufgreifen und replizieren (Bianchi et al., 2023).

Insbesondere bei unspezifischen Prompts können sich Verzerrungen in der Beschreibung und Darstellung von Ethnien, sexueller Orientierung und Behinderungen ergeben (Thomson et al., 2024) oder in rassistischen und misogynen Inhalten niederschlagen (Ananya, 2024). Dabei kann es nicht nur zur unbeabsichtigten, sondern auch missbräuchlichen Erstellung solcher Inhalte kommen. Akteur*innen können mittels Prompt Injection Sicherheitslücken von KI ausnutzen, um Regelungen zu umgehen und gefährdende Informationen oder verbotene Bildinhalte zu erstellen (Struppek et al., 2023).

Zu beachten ist weiterhin, dass TTI-Modelle für die Erstellung von Bildern auf bereits bestehende Inhalte zurückgreifen und somit Bilder generieren, die fast identisch mit der jeweiligen Quelldatei sein können. Dadurch besteht auch das Risiko von Urheberrechtsverletzungen beim Umgang mit generativer KI (Bird et al., 2023).



Rechts: Beispiel-Ergebnisse aus Midjourney für den Prompt „Taxifahrer im Taxi“.

oder Handlungen zu stehlen (*Liar's Dividend*). Zusammengefasst besteht die Gefahr, dass sich Bürger*innen manipuliert fühlen und Misstrauen gegenüber politischen Akteur*innen entwickeln, was zu generellen Zweifeln am politischen Diskurs und einem Vertrauensverlust in demokratische Institutionen führen könnte (Battista, 2024; Jungherr, 2024; Jungherr & Schroeder, 2023; Novelli & Sandri, 2024).

2.4 Regulierungsbedarf für politische KI-Kampagnen? Aktuelle Vorschläge und Diskussionen

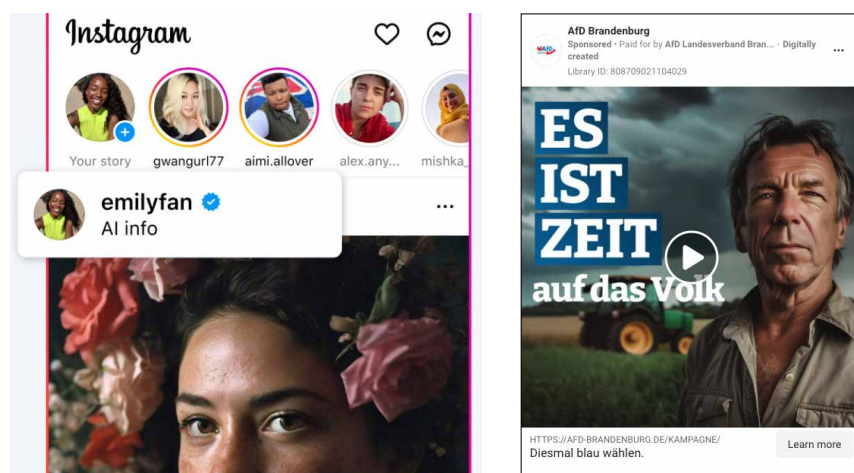
Der rasante Fortschritt und die zunehmende Verbreitung von generativer KI in der politischen Kommunikation werfen dringende Fragen zur Regulierung dieser Technologie auf (Chowdhury, 2024; G'sell, 2024; Jungherr, 2024). Die Herausforderung effektiver Regulierung besteht darin,

einen Rahmen zu schaffen, der einerseits die Innovationskraft und die potenziellen Vorteile von KI für den demokratischen Diskurs fördert, andererseits aber auch die Integrität des politischen Prozesses und das Vertrauen der Bürger*innen schützt. Gleichwohl muss jedoch die Balance zwischen dem Schutz der Demokratie und der Wahrung der freien Meinungsäußerung gewahrt werden, was eine sorgfältige Abwägung der verschiedenen Interessen erfordert (G'sell, 2024; Laude & Daum, 2024). In diesem Kapitel betrachten wir verschiedene Ansätze zur Regulierung des KI-Einsatzes in politischen Kampagnen, von freiwilligen Selbstverpflichtungen der Parteien über Transparenzmaßnahmen und unabhängige Kontrollen bis hin zu möglichen gesetzlichen Verboten.

Transparenzhinweise: Die Kennzeichnung von KI-generierten Inhalten durch Disclaimer oder

Abbildung 9:

Transparenzhinweis für die KI-Nutzung auf Instagram in unbezahlten Posts (links) und bezahlter Werbung (rechts)



Quellen: Meta, 2024 / Meta Ad Library/AfD Brandenburg.

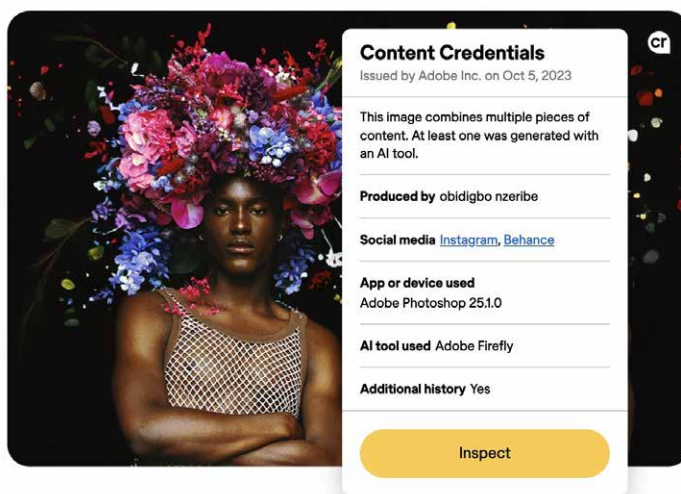
Wasserzeichen wird sowohl in Gesetzesinitiativen wie dem im März 2024 verabschiedeten europäischen „AI-Act“ (Europäisches Parlament, 2024; siehe für weitere Gesetzesinitiativen G’sell, 2024) als auch in Selbstverpflichtungsmaßnahmen der Parteien und Plattformen (siehe weiter unten) als wichtige Maßnahme zur Regulierung des KI-Einsatzes in politischen Kampagnen diskutiert. So verpflichtet der AI-Act Anbieter*innen von KI-generierten Bild-, Audio- oder Videoinhalten zur Kennzeichnung, dass die Inhalte künstlich erzeugt oder manipuliert wurden. Denn laut des AI-Acts wird KI, die zur Beeinflussung des Wahlverhaltens verwendet wird, als hochrisikant eingestuft. Jedoch existieren aktuell keine spezifischen Ablaufpläne und Mechanismen zur Kennzeichnung oder Entfernung nicht gekennzeichneteter KI-Inhalte (Laude & Daum, 2024).

Als Reaktion darauf begann Meta auf seinen Plattformen Facebook und Instagram, KI-generierte Bilder und Videos in unbezahlten Posts mit dem Label „KI-Info“ zu versehen und weitere Informationen zum Ursprung abrufbar zu machen (Meta, 2024; siehe Abbildung 9 links). Jedoch kennzeichnet Meta KI-Inhalte in bezahlter Werbung mit dem Label „Digitally Created“ (Meta 2024; siehe Abbildung 9 rechts), was ein Beleg für die aktuellen unregulierten und unsystematischen Mechanismen zur Kennzeichnung ist.

OpenAI beabsichtigt, digitale und sichtbare Wasserzeichen in seinen Bildgenerator DALL-E zu integrieren. Dabei setzt es auf das „Content-Credentials“-Hinweislabel (CR), das von der „Coalition for Content Provenance and Authenticity“ (C2PA) entwickelt wurde (OpenAI, 2024b, siehe Abbildung 10). Das Symbol wird in Bildern und Videos eingebettet und gibt beim Anklicken

Abbildung 10:

Transparenzhinweis für die KI-Nutzung durch DALL-E



Quelle: OpenAI, 2024.

Aufschluss über Quelle, Erstellungsdatum und verwendete KI-Werkzeuge. Diese Kennzeichnung garantiert jedoch ebenfalls keine wasserdichte Lösung, weil sie im Nachgang der Erstellung aus den Metadaten des Bildes entfernt werden kann. Dies möchte OpenAI jedoch ahnden und entsprechende Erstelleraccounts von nachträglich bearbeiteten Bildern auf ihren Plattformen sperren.

Die Idee hinter Transparenzhinweisen basiert auf Erkenntnissen zum Konzept des sogenannten Persuasionswissens (Friestad & Wright, 1994, 1995). Mediennutzer*innen entwickeln im Laufe ihrer Mediensozialisation ein Verständnis für persuasive Botschaften und können dementsprechend darauf reagieren. Dieses Wissen dient als kognitiver Abwehrmechanismus und kann die Verarbeitung und Wirkung von Werbeeinhalten beeinflussen (Cowley & Barron, 2008; Jost et al., 2023; Matthes et al., 2007; Nelson et al., 2021). Bei KI-generierten Inhalten ist die Aktivierung dieses Persuasionswissens besonders wichtig, da diese oft schwer von menschlich erstellten zu unterscheiden sind (u. a. Frank et al., 2023; Groh et al., 2024; Köbis et al., 2021). Eine klare Kennzeichnung kann das Persuasionswissen aktivieren und somit eine kritischere Auseinandersetzung mit den Botschaften fördern (Beckert et al., 2021; Boerman et al., 2012, 2017). Dies ist umso bedeutsamer, als bei vielen Menschen noch wenig Wissen über KI besteht, wodurch die Erkennung KI-generierter Inhalte oft allein von einer entsprechenden Kennzeichnung abhängt (für kritische Auseinandersetzung siehe Altay & Gilardi, 2024; Leibowicz, 2024).

*Unabhängige Expert*innenprüfung:* Ein weiterer Vorschlag zur Regulierung von KI in politischen Kampagnen ist die Einführung einer unabhängigen Expert*innenprüfung für KI-generierte politische Inhalte vor deren Veröffentlichung nach dem Vorbild von Faktenchecks (Bieber et al., 2024; McKay et al., 2024). Dieser Ansatz würde eine zusätzliche Kontrollinstanz schaffen, die die Integrität und Angemessenheit der KI-Inhalte bewertet und darüber auf Basis von vorhandener Expertise entscheidet. Unabhängige Expert*innen könnten potenzielle Risiken durch KI wie Desinformation, Manipulation oder die Verstärkung von Vorurteilen identifizieren und filtern (G'sell, 2024; Jungherr, 2024). Eine solche Maßnahme basiert unter anderem darauf, dass Wissenschaftler*innen Zugang zu Plattformdaten haben, um kurzfristig Ad-Hoc-Entscheidungen abzustimmen oder langfristig die Auswirkungen von KI-generierten Botschaften und die wirksamsten Reaktionen darauf besser bewerten zu können (Budak et al., 2024; Klinger & Ohme, 2023). Dies könnte dazu beitragen, das Vertrauen der Öffentlichkeit in politische Kommunikation zu stärken, birgt aber auch Herausforderungen hinsichtlich der Praktikabilität und möglicher Einschränkungen der freien Meinungsäußerung. Letzteres Argument nutzten die Social-Media-Unternehmen von Facebook, Instagram und X in der Debatte um unabhängige Faktenchecks von Inhalten und wollen den Wahrheitsgehalt von diesen nicht mehr durch ausgewählte Organisationen oder Nachrichtenagenturen überprüfen lassen (Kaplan, 2025).

Selbstverpflichtungsmaßnahmen der Parteien:

Unter selbstregulatorische Maßnahmen fällt die freiwillige Einschränkung politischer Parteien bei der Nutzung von KI in ihren Kampagnen oder ein kompletter Verzicht (G'sell, 2024; Jungherr, 2024). Dieser Ansatz basiert auf der Idee der Selbstverpflichtung und ethischen Verantwortung der politischen Akteur*innen. Bislang haben nur wenige deutsche Parteien offizielle Richtlinien zum KI-Einsatz in der politischen Arbeit veröffentlicht. Die SPD-Bundestagsfraktion verpflichtet sich in einem Positionspapier zu Transparenz, Datenschutz und Ressourcenbewusstsein bei der Nutzung von KI (SPD-Bundestagsfraktion, 2024). Auch die FDP hat eine KI-Selbstverpflichtung veröffentlicht, die insbesondere die transparente Nutzung, Informationsintegrität, den Datenschutz und ethische Verantwortung betont (FDP-Bundestagsfraktion, 2023). Die Grünen haben zwar interne Richtlinien entwickelt, die Transparenz und den Verzicht auf Stimmenimitation durch KI vorschreiben, haben diese aber noch nicht veröffentlicht (Müllender, 2024). Auf Landesebene haben bisher nur Die Linke in Sachsen und Thüringen KI-Richtlinien beschlossen (DIE LINKE. Sachsen Landesvorstand, 2023; DIE LINKE. Thüringen, 2023). CDU/CSU, AfD und BSW haben bislang keine spezifischen Positionen zum KI-Einsatz veröffentlicht. Auch die großen US-Parteien verfügen noch nicht über entsprechende Richtlinien (Merica, 2024).

Ein weiteres Beispiel für einen solchen selbstregulierenden Ansatz sind Verhaltenskodizes für Wahlen wie der zur Bundestagswahl 2025 (Spie-

gel.de, 2024) oder zum Europäischen Parlament 2024 (European Commission, 2024; siehe für weitere Verhaltenskodizes G'sell, 2024). Diese Kodizes sehen für alle unterzeichnenden Parteien einen verantwortungsvollen Umgang mit KI vor und verpflichten dazu, die Integrität der Wahlen zu wahren und einen fairen Wahlkampf zu führen. In Bezug auf den Einsatz von KI in politischen Kampagnen sehen die Kodizes vor, dass keine irreführenden Inhalte erstellt, verwendet oder verbreitet werden dürfen. Insbesondere wird der Einsatz von KI-generierten Audio-, Bild- oder Videomaterialien, die Kandidat*innen oder Amtsträger*innen in irreführender Weise darstellen oder ihnen Aussagen in den Mund legen, untersagt (European Commission, 2024; Spiegel.de, 2024). Sollten KI-generierte Inhalte verwendet werden, müssen diese eindeutig als solche gekennzeichnet werden. Der Vorteil von Selbstverpflichtungsmaßnahmen liegt in ihrer Flexibilität und der Möglichkeit, schnell auf neue Entwicklungen zu reagieren. Allerdings hängt die Wirksamkeit stark von der Bereitschaft aller Parteien ab, sich an solche Vereinbarungen zu halten, und es fehlen oft durchsetzbare Sanktionsmechanismen bei Verstößen.

Verbot von KI-generierten Inhalten: Die restriktivste Form der Regulierung wäre ein vollständiges Verbot von KI-generierten Inhalten in politischen Kampagnen. Dieser Ansatz würde jegliche Verwendung von KI zur Erstellung politischer Botschaften untersagen, um potenzielle Risiken wie Manipulation und Desinformation von vornherein auszuschließen. Befürworter*in-

nen argumentieren, dass ein solches Verbot notwendig sein könnte, um die Integrität des demokratischen Prozesses zu schützen, insbesondere angesichts der schnellen Entwicklung und Missbrauchsmöglichkeiten von KI-Technologien (Bundesrat, 2024; G'sell, 2024; Laude & Daum, 2024). So ist es beispielsweise politischen Parteien und Kandidat*innen in Brasilien laut Wahlrecht verboten, KI-generierte Inhalte zu nutzen, um einer Kandidatur zu schaden oder sie zu fördern, einschließlich der Verwendung zur „Schaffung, Ersetzung oder Veränderung des Bildes oder der Stimme einer lebenden, verstorbenen oder fiktiven Person“ (Tardáguila, 2024).

Anbieter von KI-Systemen reagieren auf eine solche Forderung, indem sie die Erstellung bestimmter Inhalte einschränken. So ist zum Beispiel die Erstellung von Texten für strafrechtliche Sachverhalte wie terroristische Aktivitäten und Verleumdungen oder Bilder beziehungsweise Videos von Politiker*innen oft nicht möglich, da die Anbieter Sperren in die Systeme programmiert haben (OpenAI, 2024a, 2024b; Wankhede, 2024). Es gibt jedoch auch KI-Tools ohne solche Begrenzungen, wie beispielsweise „Grok“ von X, die kaum einer Kontrolle unterliegen und dadurch die Verbreitung von extremen Inhalten ermöglichen (siehe Abbildung 2 rechts; Robertson, 2024).

Im deutschen Kontext diskutieren Rechtsexpert*innen, ob der Einsatz von KI-generierten Inhalten in politischen Kampagnen gegen Artikel 21 Abs. 1 GG verstoßen könnte, da Parteien zur Förderung des demokratischen Willensbil-

dungsprozesses verpflichtet sind, richtige Informationen zu verbreiten (Laude & Daum, 2024). Auch ein vom Freistaat Bayern in den Bundesrat eingebrachter Gesetzesentwurf zielt darauf ab, einen spezifischen Straftatbestand im Strafgesetzbuch zu schaffen, um den Missbrauch von KI-generierten Medieninhalten zu bekämpfen, die den Anschein authentischer Aufnahmen – insbesondere von Menschen – erwecken (Bundesrat, 2024). Kritiker*innen weisen jedoch darauf hin, dass ein teilweises oder vollständiges Verbot schwer durchsetzbar wäre, weil es zu stark das Recht auf freie Meinungsäußerung einschränke (Chowdhury, 2024; Jungherr, 2024; Łabuz & Nehring, 2024). Außerdem verhin-dere es möglicherweise legitime und innovative Nutzungen von KI in der politischen Kommunikation.

2.5 Zwischenfazit: Der kontroverse Beginn einer technologischen Revolution

Generative KI hat sich innerhalb kurzer Zeit zu einem wichtigen Instrument der politischen Kampagnenkommunikation entwickelt. Die in diesem Kapitel diskutierten Beispiele zeigen, dass die Technologie bereits vielfältig eingesetzt wird – von der Erstellung von Wahlkampfmaterialien über die Personalisierung von Botschaften bis hin zur gezielten Desinformation. Dabei werden sowohl die Potenziale als auch die Risiken dieser Entwicklung deutlich: Einerseits ermöglicht KI eine effizientere und inklusivere politische Kommunikation, die auch ressourcenschwächeren Akteur*innen professionelle Kampagnen ermöglicht. Andererseits birgt sie erhebliche Risi-

ken für die demokratische Willensbildung, etwa durch die massenhafte Verbreitung von Falschinformationen oder den möglichen Vertrauensverlust in politische Kommunikation.

Diese Ambivalenz spiegelt sich auch in den aktuellen Regulierungsansätzen und -vorschlägen wider. Während einige Parteien durch Selbstverpflichtungen und Transparenzrichtlinien einen verantwortungsvollen Umgang mit KI anstreben, fehlen bei anderen jegliche Regelungen. Auch auf Seiten der Plattformen, KI-Unternehmen und Regulator*innen wird zwischen notwendigem Schutz demokratischer Prozesse und der Wahrung der Meinungsfreiheit abgewogen. Diese Debatte intensivierte sich mit der Entscheidung der Social-Media-Unternehmen von Facebook, Instagram und X, künftig KI zum festen Bestandteil bei der Erstellung von Botschaften zu machen (Robertson, 2024; Sato, 2025; Wankhede, 2024) und den Wahrheitsgehalt von Inhalten nicht mehr durch unabhängige Organisationen oder Nachrichtenagenturen überprüfen lassen zu wollen (Kaplan, 2025).

Vor diesem Hintergrund ist es essenziell zu verstehen, wie die Bevölkerung den Einsatz von KI in der politischen Kommunikation wahrnimmt und welche Wirkungen sie bei den Bürger*innen entfaltet. Denn KI-Kampagnen können zu ganz unterschiedlichen Wirkungen führen, je nachdem, welcher Kenntnisstand über die Technologie vorhanden ist, ob sie Emotionen wie Angst oder Hoffnung auslöst, als glaubwürdig bewertet oder mit Potenzialen oder Risiken verbunden wird.

In den beiden folgenden empirischen Untersuchungen erforschen wir daher erstmals systematisch die Einstellungen der deutschen Bevölkerung zu KI in politischen Kampagnen sowie die Wahrnehmung und Wirkung KI-generierter Kampagneninhalte. Ziel ist es, mit den Ergebnissen ein Bewusstsein für die Potenziale und Risiken des Einsatzes dieser Technologie in der politischen Kommunikation zu schaffen und damit einen informierten öffentlichen Diskurs über die verantwortungsvolle Nutzung durch politische Akteur*innen anzuregen.

3 Befragung zum Einsatz von generativer KI in politischen Kampagnen

Während bereits verschiedene Studien allgemeine Einstellungen der Bevölkerung zu KI oder deren Einsatz in wirtschaftlichen Bereichen untersucht haben, ist über die Wahrnehmung und Bewertung von KI in der politischen Kommunikation aktuell noch nichts bekannt. Diese Forschungslücke wiegt besonders schwer angesichts der zunehmenden Bedeutung von KI in politischen Kampagnen und daraus resultierenden gesellschaftlichen Debatten (siehe Kapitel 2). Wir möchten zur Schließung dieser Lücke mit einer repräsentativen Befragung der deutschen Bevölkerung über den Einsatz generativer KI in politischen Kampagnen beitragen.

Im Folgenden führen wir zunächst weiter in unser Erkenntnisinteresse ein, stellen die konkreten Forschungsfragen vor und ordnen diese in den aktuellen Stand der Forschung ein (3.1). Im Anschluss erläutern wir Methode und Aufbau der Befragung (3.2) und diskutieren dann die Ergebnisse (3.3). Das Kapitel endet mit einem Zwischenfazit, in dem die wichtigsten Ergebnisse für die einzelnen Forschungsfragen noch einmal zusammengefasst werden (3.4).

3.1 Erkenntnisinteresse und Forschungsfragen

Die wissenschaftliche Auseinandersetzung mit den Auswirkungen generativer KI auf gesell-

schaftliche Prozesse hat in den letzten Jahren deutlich zugenommen. Forschende untersuchen dabei insbesondere, wie Menschen die Chancen und Risiken dieser Technologie wahrnehmen, wie sie Organisationen oder Akteur*innen bewerten, die KI einsetzen, und unter welchen Bedingungen sie KI-Anwendungen akzeptieren. Für den Kontext der politischen Kommunikation fehlen jedoch noch Erkenntnisse. Im Folgenden stellen wir unsere sechs übergeordneten Forschungsfragen vor, die unsere repräsentative Befragung leiten. Dabei ordnen wir jede Frage in den aktuellen Forschungsstand ein und erläutern unser Erkenntnisinteresse.

Fragestellung 1: Welche Kenntnisse haben Bürger*innen über KI in politischen Kampagnen?

Eine fundierte Medienkompetenz ist entscheidend für eine informierte und kritische Auseinandersetzung mit politischen Botschaften (Hugger, 2022; Iske & Barberi, 2022; Wang et al., 2024). Eine solche Kompetenz muss inzwischen auch Kenntnisse über die Funktionsweisen, Einsatzfelder und Nutzungsmöglichkeiten von (generativer) KI einschließen. Solche Kenntnisse stärken nicht nur die gesellschaftliche Resilienz gegenüber technologischen Veränderungen (Sundar, 2008), sondern ermöglichen es den Bürger*innen auch, politische Kommunikationsabläufe besser zu verstehen und einzuordnen (Bimber, 2003; Dommert et al., 2023;

Wang et al., 2024). Dies gilt insbesondere für die Erkennung von KI-Inhalten, die nur auf Basis eines vorhandenen Wissens über generative KI entlarvt werden können (Frank et al., 2023; Groh et al., 2024; Köbis et al., 2021), um sich vor allem vor einer möglichen Nutzung von KI für Halb- oder Unwahrheiten zu schützen. Aber welche Erfahrungen haben Bürger*innen bisher mit dem Einsatz von KI in politischen Kampagnen gemacht? Und welches Verständnis haben sie von diesem Einsatz?

Eine repräsentative Bevölkerungsumfrage des Bayerischen Forschungsinstituts für Digitale Transformation zum allgemeinen Wissensstand über KI zeigt, dass zwar eine Mehrheit von generativer KI gehört hat, aber nur ein kleiner Teil tatsächliche Erfahrungen damit gemacht hat. Gerade einmal ein Drittel der Befragten hat KI-Systeme schon mindestens einmal genutzt (Schlude et al., 2023). Eine repräsentative Onlinestudie von ARD und ZDF zeigt zudem auf, dass der Begriff der generativen KI über Altersgruppen hinweg große Bekanntheit aufweist (Beisch & Koch, 2023). Eine tiefere Kenntnis bedeutet dies allerdings nicht. Eine Erhebung des Meinungsmonitors Künstliche Intelligenz zeigt, dass nur jede fünfte Person ihr Wissen als (eher) hoch einschätzt, während mehr als ein Viertel angibt (eher) nichts zu wissen. Dabei gibt mehr als ein Drittel an, KI bereits im Alltag zu nutzen (MeMo:KI, 2022). Allerdings wird generative KI vornehmlich von unter 30-Jährigen genutzt, während ältere Personen zwar Kenntnisse über diese Technologie besitzen, sie aber nicht verwenden (Beisch & Koch, 2023). Mit Blick auf das Wissen

über Deepfakes zeigt eine repräsentative Befragung in Deutschland aus dem Jahr 2022 ähnliche Befunde (Bitton et al., 2024): Rund zwei Drittel der Menschen haben keinerlei Vorkenntnisse bei dem Thema und rund drei Viertel der Befragten glaubten, bisher noch keinen bewussten Kontakt mit Deepfakes gehabt zu haben.

Diese Ergebnisse zu den allgemeinen Kenntnissen über generative KI und Deepfakes unterstreichen die Notwendigkeit, den aktuellen Wissensstand der Bevölkerung über den Einsatz von KI im Kontext von politischen Kampagnen zu erfassen.

Fragestellung 2: Wie beurteilen Bürger*innen den Einsatz generativer KI in politischen Kampagnen?

Die allgemeine Bewertung des Einsatzes von KI in politischen Kampagnen hat große Bedeutung für die Legitimität politischer Kommunikation und die strategischen Nutzungsmöglichkeiten durch politische Akteur*innen. Wenn die Bevölkerung den Einsatz von KI als (nicht) wünschenswert, sicher oder gefährlich, (un)fair oder (in) akzeptabel empfindet, kann sich das auf das Vertrauen der Bürger*innen in demokratische Institutionen oder die Legitimität von Wahlen auswirken (Jungherr, 2023; Jungherr & Schroeder, 2023). Außerdem beeinflusst die Bewertung von KI die strategischen Möglichkeiten und ethischen Standards, die für den Einsatz von KI festgelegt und befolgt werden, weil politische Akteur*innen die Sorgen der Bevölkerung berücksichtigen müssen (Novelli & Sandri, 2024). Darüber hinaus spielt die öffentliche Meinung zu KI eine entscheidende Rolle bei der möglichen

Ausgestaltung von Regularien und Gesetzen. Denn eine Bevölkerung, die KI in der Politik als bedrohlich oder gefährlich einstuft, könnte strengere Regulierungen oder Einschränkungen für deren Einsatz fordern (G'sell, 2024; Jungherr & Rauchfleisch, 2024).

Bisherige Studien zur allgemeinen Akzeptanz von KI haben sich auf Bereiche außerhalb der Politik konzentriert. Hier zeigt sich eine durchaus gesplante Meinung: Eine repräsentative Befragung der deutschsprachigen wahlberechtigten Bevölkerung im Auftrag der Konrad Adenauer Stiftung ergibt, dass 24 Prozent der Befragten KI als Ursache neuer Probleme wahrnehmen, wohingegen knapp 29 Prozent KI vorwiegend als Lösung bestehender Probleme sehen (Neu, 2024).

Eine Befragungsstudie mit Gruppendiskussion im Auftrag der Bayerischen Landeszentrale für neue Medien zeigt, dass auch unter Jugendlichen sowohl positive als auch negative Reaktionen auf KI zu finden sind (Wendt et al., 2024). Einerseits bewerteten sie KI überwiegend positiv und zeigten sich begeistert von den kreativen Möglichkeiten der Technologie. Andererseits äußerten sie auch Bedenken – insbesondere aufgrund von Unsicherheiten über urheberrechtliche Fragen und mögliche zukünftige Entwicklungen der Technologie.

Auch in Deepfakes sehen fast alle Befragte einer repräsentativen Befragung in Deutschland eine hohe (69,9 Prozent) beziehungsweise mittlere (24,3 Prozent) Gefahr (Bitton et al., 2024). Nur knapp 6 Prozent schätzen die mit Deepfakes ver-

bundenen Risiken als niedrig ein. Je jünger und gebildeter eine Person ist, desto eher hat sie auch eine Vorstellung von Chancen und (positiven) Einsatzfeldern für Deepfakes. Vorteile von Deepfakes sehen die Befragten für die Videospiel-, Mode- und Kunstbranche.

Um den aktuellen Kenntnisstand auf den noch nicht untersuchten Kontext von politischen Kampagnen zu übertragen, erfragen wir in einem zweiten Schritt die allgemeine Wahrnehmung und Beurteilung von generativer KI in politischen Kampagnen.

Fragestellung 3: Wie bewerten Bürger*innen Parteien, die generative KI einsetzen?

Die Parteibewertung spielt eine wesentliche Rolle für das Vertrauen der Wähler*innen in eine Partei und kann deren Wählbarkeit beeinflussen (Maurer, 2014; Schmitt-Beck, 2000). In diese Bewertung kann auch der Einsatz von Wahlkampftechnologien wie KI einfließen (Dommett et al., 2023; Jungherr et al., 2020). So können Parteien, die generative KI einsetzen, von Bürger*innen einerseits als innovativ und zukunftsorientiert wahrgenommen werden, andererseits aber auch Bedenken hinsichtlich Authentizität und Vertrauenswürdigkeit auslösen – insbesondere, wenn der Einsatz von KI als Versuch zur Manipulation oder Täuschung der Wähler*innen wahrgenommen wird (Jungherr, 2016; Novelli & Sandri, 2024; Sandri et al., 2024).

Bisherige Umfragen ohne Bezug auf politische Kampagnenarbeit zeigen, dass die Akzeptanz

von KI-Anwendungen stark vom Kontext abhängt: Bei interpersoneller Kommunikation im Dienstleistungsbereich zeigen Jakesch et al. (2019), dass potenzielle Kund*innen Anbieter*innen als weniger vertrauensvoll wahrnehmen, wenn sie befürchten, dass Texte auf den Profilen der Anbieter*innen KI-generiert sein könnten. Im Bereich des Journalismus zeigt eine Befragung der Schweizer Bevölkerung von Vogler et al. (2023) eine sehr deutliche Präferenz für Medienbeiträge, die ohne KI erstellt wurden und legt nahe, dass die Arbeit von und das Vertrauen in Journalist*innen nach wie vor sehr wichtig ist.

In diesem Kontext sind auch die Ergebnisse einer repräsentativen Befragung deutscher Bürger*innen über Deepfakes interessant (Bitton et al., 2024): Diese legen nahe, dass das Aufkommen von Deepfakes insgesamt dafür sorgt, dass Menschen Medien weniger Vertrauen entgegenbringen – und zwar unabhängig davon, ob sie tatsächlich selbst mit Deepfakes in Berührung kommen oder nicht. Insbesondere Befragte, die ihre Fähigkeiten, Deepfakes erkennen zu können, höher einschätzen, entwickeln dabei ein Gefühl von Medienzynismus und machen sich mehr Sorgen über den Missbrauch durch politische oder mediale Institutionen.

Die bisherigen Ergebnisse aus den nicht-politischen Kontexten verdeutlichen die Notwendigkeit zu erforschen, wie Bürger*innen Parteien bewerten, die KI einsetzen. Entsprechend zielt unser nächstes Erkenntnisinteresse auf Fragen des Vertrauens in und der Wählbarkeit von Parteien ab, die KI einsetzen.

Fragestellung 4: Inwieweit akzeptieren Bürger*innen den Einsatz generativer KI für bestimmte Aufgaben bei der Wählerkommunikation?

Die Akzeptanz der Nutzung von KI für bestimmte Kampagnenaufgaben ist maßgeblich für das Vertrauen in die Parteikommunikation sowie deren Legitimität und Effektivität. Einerseits können einige Aufgaben durch den unterstützenden Einsatz von KI, wie etwa bei sprachlichen Korrekturen oder Übersetzungen, das Vertrauen in die Kompetenz einer Partei stärken, weil sie die Professionalität der Kommunikation potenziell erhöhen, ohne inhaltliche Aspekte zu beeinflussen (Jungherr et al., 2020; Novelli & Sandri, 2024). Andererseits könnten Aufgaben, bei denen KI stärker in die inhaltliche Gestaltung eingreift oder diese vollständig übernimmt, erhebliche Risiken für einen Vertrauensverlust und Legitimitätsprobleme bergen (Battista, 2024; Novelli & Sandri, 2024). Dies liegt daran, dass in politischen Kampagnen oft wichtige politische Positionen und Entscheidungen kommuniziert werden, sodass die Authentizität und Verantwortlichkeit der Parteienkommunikation auf dem Spiel stehen. Letztlich besteht besonders bei Aufgaben, die die direkte Kommunikation mit Wähler*innen betreffen, die Gefahr, dass die KI-generierte Kommunikation als potenziell manipulativ oder unaufrichtig wahrgenommen wird (Dobber et al., 2021; Łabuz & Nehring, 2024; Sætra, 2023; Vaccari & Chadwick, 2020). Dies kann das Vertrauen in die und die Legitimität der Parteikommunikation stark beeinträchtigen (Jungherr & Schroeder, 2023).

Bisherige Studien zeigen, dass die Akzeptanz von KI je nach Anwendungsbereich und Komplexität der Aufgaben stark variiert. Beispielsweise stehen Bürger*innen dem Einsatz von KI im Journalismus skeptisch gegenüber, insbesondere wenn es um vollständig KI-generierte Inhalte geht (Vogler et al., 2023; Behre et al., 2024). Höhere Akzeptanz eines KI-Einsatzes wird bei weniger komplexen Aufgaben wie der Routineberichterstattung zum Wetter oder zu Börsenkursen beobachtet oder für die Berichterstattung im Softnews-Bereich, das heißt über Stars und Celebrities und Sport. Dagegen ist die Zustimmung in den Bereichen von Politik, Wirtschaft und Wissenschaft deutlich niedriger. Besonders unter älteren Bevölkerungsgruppen besteht ein starkes Misstrauen gegenüber dem Einsatz von KI für unterschiedliche Aufgaben etwa im Journalismus oder in Kunst und Kultur (Kero et al., 2023).

Die Akzeptanz von KI variiert laut repräsentativen Befragungen in den Niederlanden auch bei Aufgaben in der Gesellschaftspolitik (Votta & de León, 2024): Beispielsweise findet knapp mehr als die Hälfte der Befragten, dass die Verteilung von Sozialstaatsmitteln durch KI nicht sinnvoll ist. Dennoch sehen einige Befragte die Nützlichkeit, die KI unter anderem im Bereich des Gesundheitssystems oder in der Strafverfolgung schwerer Verbrechen bietet.

Aufbauend auf dem aktuellen Forschungsstand in anderen gesellschaftlichen Bereichen, erheben wir mit der vierten Forschungsfrage, welche Anwendungen von KI in politischen Kampagnen

von der Bevölkerung als akzeptabel oder inakzeptabel angesehen werden. Dabei unterscheiden wir zwischen solchen Aufgaben, bei denen generative KI lediglich unterstützend eingesetzt wird, und solchen, bei denen generative KI zur Erstellung von Inhalten eingesetzt wird.

Fragestellung 5: Welche Chancen und Risiken sehen Bürger*innen beim Einsatz generativer KI in politischen Kampagnen?

Die Auseinandersetzung mit den Chancen und Risiken des Einsatzes von generativer KI in politischen Kampagnen ist essenziell für demokratische Wahlen und Diskurse. Dies haben wir ausführlich in Kapitel 2 besprochen und gezeigt, dass diese Technologie einerseits effizientere und inklusivere Kampagnen ermöglichen, die auch weniger finanzstarken Akteur*innen zugutekommen und durch personalisierte Inhalte ausgewählte Teilöffentlichkeiten einschließen können; auf der anderen Seite können ein hoher Ressourcenverbrauch, mögliche Urheberrechtsverletzungen, die Gefahr von Desinformation oder die Verstärkung von Vorurteilen und Misstrauen in den politischen Diskurs erhebliche Probleme darstellen.

Bisherige Befragungen zeigen eine eher negative Einschätzung der Bevölkerung, die KI als potenzielle Gefahr für die Demokratie betrachtet (Bitton et al., 2024; Bühler, 2023; Fletcher & Nielsen, 2024; Kieslich et al., 2022; Vogler et al., 2023; Votta & de León, 2024). Dabei befürchten die Befragten insbesondere eine beschleunigte Verbreitung von Desinformationen und staatlicher Propaganda (Bühler, 2023; Kieslich et al.,

2022), aber auch eine Verschlechterung der Qualität des Journalismus (Vogler et al., 2023).

Länderübergreifende Studien deuten darauf hin, dass die Mehrheit der Befragten negative Auswirkungen auf politische Parteien, Nachrichtenmedien, Wissenschaft und Gesellschaft erwartet (Fletcher & Nielsen, 2024). Sie befürchten unter anderem Arbeitsplatzverluste und steigende Lebenshaltungskosten, aber auch, dass KI von manchen Akteur*innen missbräuchlich verwendet werden könnte. Darüber hinaus werden Diskriminierungsrisiken durch KI besonders in wirtschaftlichen Bereichen wie Kreditvergaben und individueller Preisgestaltung genannt (Kieslich et al., 2020). Auch mit Blick auf Deepfakes haben deutsche Befragte mehrheitlich Ängste davor, dass ihre Zahl ansteigt, dass sie die öffentliche Meinung beeinflussen, dass die Grenze zwischen Realität und Fiktion verschwimmt und dass die Demokratie generell durch sie Schaden nimmt (Bitton et al., 2024). Dagegen stimmt nur weniger als die Hälfte der Menschen in Deutschland der Aussage zu, dass Deepfakes auch für sinnvolle Zwecke genutzt werden können.

Als Chance wird der Einsatz von KI von der Mehrheit der Befragten in mehreren Ländern im Bereich der Wissenschaft, Medizin und Wirtschaft wahrgenommen (Fletcher & Nielsen, 2024), beispielsweise als Treiber von Innovation (Neu, 2024) und industrieller Produktion (Kieslich et al., 2022). Laut einer Befragung von Schlude et al. (2023) wird generative KI insbesondere von ihren Nutzer*innen als positiv empfunden und

führt zu Zeitersparnissen sowie einer positiven Entwicklung ihrer Arbeitsergebnisse.

Ausgehend von dem aktuellen Forschungsstand, zielt die Untersuchung in diesem Schritt darauf ab, herauszufinden, welche Chancen oder Gefahren die Bevölkerung im Einsatz von KI in politischen Kampagnen sieht.

Fragestellung 6: Wie bewerten Bürger*innen verschiedene Formen der Regulierung von generativer KI in politischen Kampagnen?

Für eine verantwortungsvolle Kampagnenarbeit, aber auch zum Schutz demokratischer Prozesse, ist eine Regulierung des Einsatzes generativer KI in der politischen Kommunikation von zentraler Bedeutung (Chowdhury, 2024; G'sell, 2024; Jungherr, 2024). Entscheidend ist hier, Regelungen zu schaffen, die sowohl die Vorteile der Technologie nutzen als auch vor ihrem Missbrauch schützen (Europäisches Parlament, 2024). Durch klare Regeln und Transparenzmaßnahmen, wie die Kennzeichnung von KI-generierten Inhalten, kann das Vertrauen der Bürger*innen in die politische Kommunikation gestärkt und den Wähler*innen ermöglicht werden, fundierter bewerten zu können, welche Inhalte glaubwürdig und relevant für ihre (Wahl)Präferenzen sind (Friestad & Wright, 1994, 1995). Zudem können Selbstverpflichtungsmaßnahmen der Parteien oder unabhängige Expert*innenprüfungen von KI-Inhalten dazu beitragen, die Verbreitung von Desinformation und manipulativen Inhalten einzudämmen (Bieber et al., 2024; McKay et al., 2024). Insgesamt muss jedoch die Balance zwischen dem Schutz der Demokratie und der

Wahrung der freien Meinungsäußerung gewahrt werden, was eine sorgfältige Abwägung der verschiedenen Interessen erfordert (G'sell, 2024; Laude & Daum, 2024).

Aktuelle Befragungen zeigen, dass eine Mehrheit der Bürger*innen in unterschiedlichen Ländern eine stärkere Regulierung von generativer KI befürwortet, wobei die Forderung nach Kennzeichnungspflichten und Verhaltenskodizes für KI-generierte Inhalte besonders ausgeprägt ist (Bitton et al., 2024; Fletcher & Nielsen, 2024; Kieslich et al., 2020; Votta & de León, 2024). So herrscht beispielsweise in der Schweizer Bevölkerung weitgehend Konsens, dass KI-generierte oder KI-unterstützte journalistische Inhalte von den Medien als solche transparent deklariert werden sollten (Vogler et al., 2023). Auch sieht eine Mehrheit der deutschen Bevölkerung eine Steigerung der KI-Kompetenz als notwendig an. Ein allgemeines Verbot von KI-Nutzung findet jedoch weniger Zustimmung unter Befragten (Kieslich et al., 2022). Unklarheit besteht darüber, wer die Verantwortung für den Umgang mit KI-Risiken tragen sollte. Eine Mehrheit sieht hier den Gesetzgeber in der Verantwortung, einen rechtlichen Rahmen zum sicheren Einsatz von KI zu schaffen (Bühler, 2023).

Wir wollen diesen Forschungsstand erweitern und mit unserer sechsten Forschungsfrage Antworten darauf finden, welche Maßnahmen von der deutschen Bevölkerung zur Regulierung von generativer KI in politischen Kampagnen als sinnvoll erachtet werden, um einen verantwortungsvollen Einsatz von KI in politischen Kampagnen zu gewährleisten.

3.2 Methode und Aufbau der Befragung

Die Forschungsfragen beantworten wir mittels einer repräsentativen Befragung der deutschen Bevölkerung. Im Folgenden erläutern wir das methodische Vorgehen und den Aufbau der Befragung. Dabei gehen wir zunächst auf die Rekrutierung der Teilnehmenden, die Zusammensetzung der Stichprobe und anschließend auf den Aufbau und die Struktur des verwendeten Fragebogens ein.

3.2.1 Rekrutierung und Stichprobe

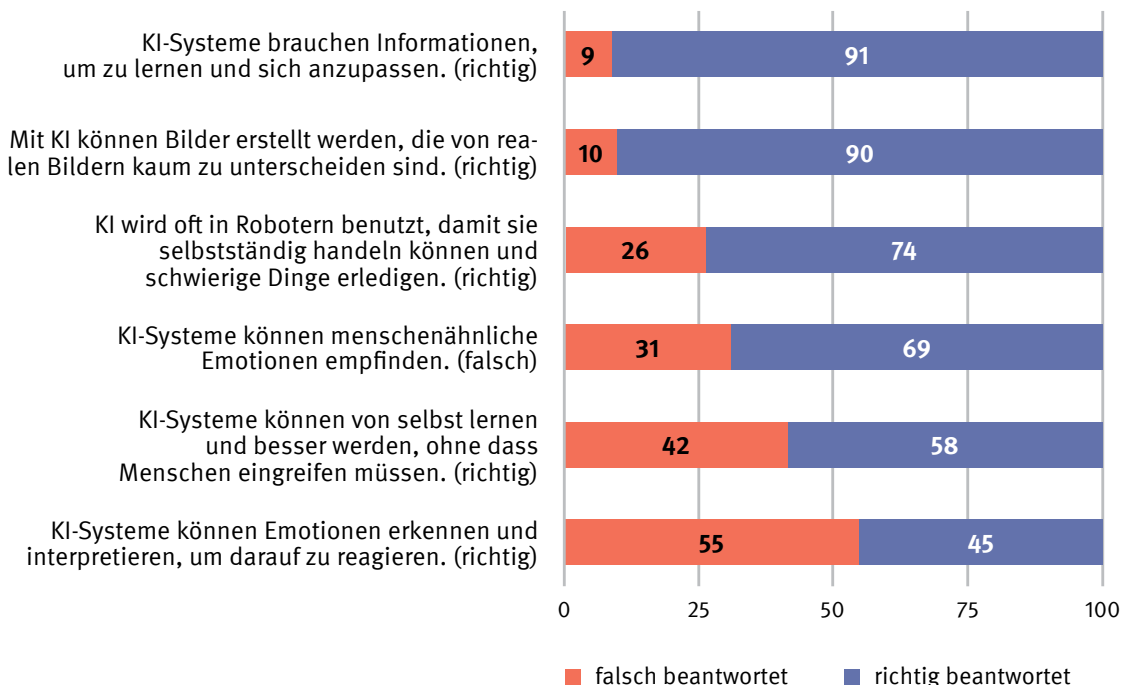
Die Rekrutierung der Teilnehmenden wurde durch den Panelanbieter Cint vorgenommen. Die befragten Personen wurden für ihre Teilnahme durch den Anbieter incentiviert. Bei der Rekrutierung wurde ein für Alter und Geschlecht (kreuzklassifiziert) bevölkerungsrepräsentatives Abbild der deutschen Bevölkerung (online) zwischen 18 und 69 Jahre avisiert. Insgesamt nahmen 3.201 Personen an der Befragung im Zeitraum vom 7. bis 25. Mai 2024 teil.

Diese Stichprobe wurde anschließend durch Qualitätschecks (Straightlining, Erinnerungsfragen) bereinigt. Nach der Bereinigung wurden für die Befragung zum Einsatz von KI 1.991 Teilnehmende berücksichtigt. Um Abweichungen von der angestrebten Quotierung zu korrigieren, wurde eine Gewichtung der Fälle vorgenommen, wobei weibliche Teilnehmerinnen aus der Altersgruppe zwischen 18 bis 29 Jahren mit dem Faktor 0.55 und männliche Teilnehmer über 60 Jahre mit dem Faktor 0.77 gewichtet wurden.

Für die gewichtete Stichprobe ergibt sich ein Anteil von weiblichen Befragten von rund 51 Prozent sowie rund 48 Prozent männliche Befragte; die übrigen Befragten identifizierten sich als „divers“. Das Durchschnittsalter der Befragten lag bei 44 Jahren, wobei je rund ein Viertel der Befragten den Altersgruppen zwischen 30 und 39, 40 und 49, sowie 50 und 59 Jahren entfielen. Die Altersgruppe 18 bis 29 Jahre machte rund 11 Prozent, die Altersgruppe zwischen 60 und 69 Jahren rund 12 Prozent aus. Mit Bezug auf die Bildung weist ein Drittel der Befragten einen Hochschulabschluss vor, während knapp ein weiteres Drittel Abi-

tur oder die Fachhochschulreife besitzt. Gut ein Viertel der Befragten verfügt über einen Realschulabschluss und etwa 9 Prozent haben einen Hauptschulabschluss. Nur ein sehr kleiner Anteil hat die Schule ohne Abschluss verlassen oder ist noch Schüler*in. In Bezug auf die berufliche Situation zeigt sich, dass etwas mehr als zwei Drittel der Befragten angestellt ist, jeweils rund 5 Prozent sind selbstständig beschäftigt, verbeamtet oder arbeitssuchend. Schüler*innen, Auszubildende und Studierende machen zusammen ebenfalls 5 Prozent der Stichprobe aus. Weitere 9 Prozent gaben an, in einer anderen beruflichen Situation zu sein.

Abbildung 11:
Allgemeines Wissen zu KI



*Frage*text: „Aktuell wird ja viel über den Einsatz und die Möglichkeiten von künstlicher Intelligenz (KI) geredet. Im Folgenden wollen wir Ihnen dazu einige Fragen stellen. Welche Aussagen über künstliche Intelligenz sind richtig, welche falsch?“
Quelle: Eigene Darstellung.

Darüber hinaus verfügen die Befragten über solide Kenntnisse über KI. Insgesamt 71 Prozent aller Fragen beantworteten die Teilnehmenden korrekt (Abbildung 11). Ein Großteil weiß, dass KI-Systeme auf Basis von Informationen trainiert werden (91%) und KI Bilder erstellen kann, die von realen Bildern kaum zu unterscheiden sind (90%). Bei komplizierteren Fragen bestehen jedoch bei einigen Befragten Wissenslücken. Mehr als ein Viertel wusste nicht, dass Roboter mithilfe von KI selbstständig handeln können (26 %) und KI-Systeme keine menschenähnlichen Emotionen empfinden können (31%). Noch weniger Befragte hatten Kenntnisse davon, dass KI selbstständig lernen kann (42%). Besonders auffällig ist, dass mehr als die Hälfte der Befragten nicht korrekt einschätzte, dass KI-Systeme Emotionen erkennen und interpretieren können (55%).

3.2.2 Aufbau der Befragung

Der Fragebogen, den die Teilnehmenden beantworteten, bestand aus mehreren Frageblöcken, die unterschiedliche Themenbereiche abdeckten (siehe Abbildung 12). Der erste Block widmete sich den soziodemografischen Angaben, um grundlegende Informationen über die Befragten zu erfassen. Im zweiten Block wurden politische Einstellungen abgefragt, darunter die Bedeutung politischer Themen, das Interesse an Politik, das politische Wissen und die Teilnahme an politischen Prozessen.

Der dritte Frageblock erfasste die Mediennutzung und untersuchte zudem, wie die Befragten Massenmedien und soziale Medien nutzten, um sich politisch zu informieren. Der vierte Block

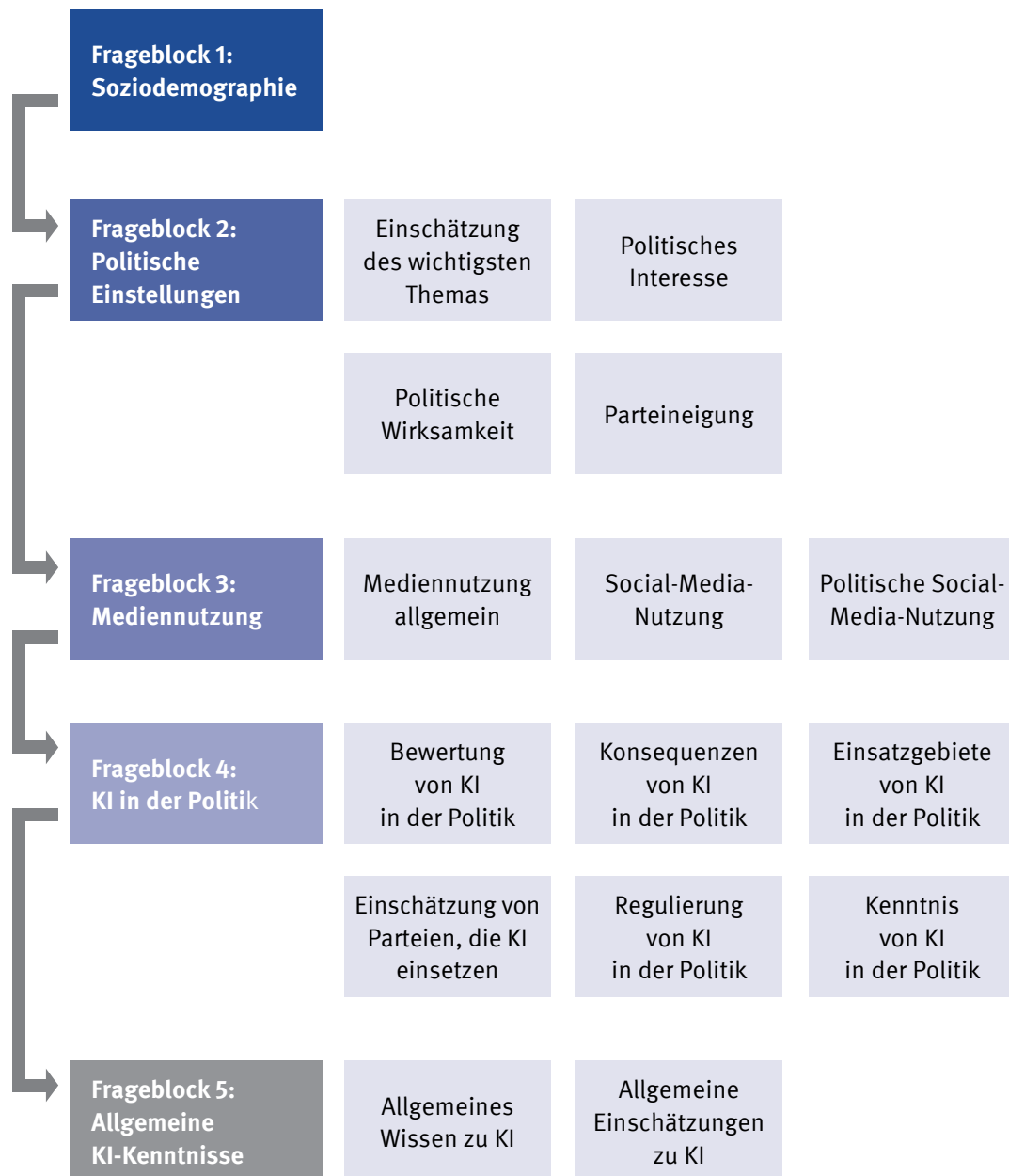
thematisierte den Einsatz von KI in der Politik, wobei die Einschätzungen der Befragten zur Bewertung von KI in der Politik, zu möglichen Konsequenzen, zu Chancen und Risiken sowie zur Regulierung von KI abgefragt wurden. Der abschließende Frageblock beschäftigte sich mit den allgemeinen Kenntnissen und Einstellungen der Befragten gegenüber KI, einschließlich Wissensfragen zur Einschätzung der KI-Kompetenz.

3.3 Ergebnisse der Befragung

Die Diskussion der Ergebnisse unserer ersten empirischen Untersuchung umfasst verschiedene Aspekte zur Bewertung des Einsatzes von KI in politischen Kampagnen und strukturiert sich nach unseren Forschungsfragen. Einleitend werden die Kenntnisse der Befragten zum Einsatz von KI in politischen Kampagnen dargestellt. Anschließend diskutieren wir, wie die Befragten den Einsatz von KI in der politischen Kommunikation beurteilen, wie sie Parteien bewerten, die KI einsetzen, und welche Einsatzgebiete von KI sie für akzeptabel halten. Abschließend diskutieren wir, welche Chancen und Risiken eines KI-Einsatzes in Kampagnen die Bevölkerung sieht und was sie von unterschiedlichen Regulierungsvorschlägen hält.

Bei der schriftlichen Ergebnisbeschreibung fassen wir jeweils die beiden äußeren Pole der Antwortmöglichkeiten zusammen (z. B. „stimme überhaupt nicht/eher nicht zu“; „stimme voll und ganz/eher zu“), um klare Tendenzen in den Antworten der Befragten zu zeigen.

Abbildung 12:
Aufbau des Fragebogens



Quelle: Eigene Darstellung.

3.3.1 Kenntnisse vom KI-Einsatz

Ein großer Teil der Befragten scheint wenig Berührungspunkte mit KI-generierten Inhalten in politischen Kampagnen zu haben (siehe Abbildung 13). Rund 60 Prozent der Befragten haben noch keine Wahlwerbung bewusst wahrgenommen, in der Text oder Bilder von KI erzeugt wurden. Zudem ist etwa der Hälfte der Befragten (eher) nicht bewusst, dass Parteien KI verwenden, um Wahlwerbung zu erstellen. Dennoch stimmen nur wenige der Aussage „eher“ oder „voll und ganz“ zu, dass Parteien fast alle Texte und Bilder mit KI erstellen (16 %). Insgesamt deutet dies auf ein gewisses Bewusstsein für den Einsatz von KI in

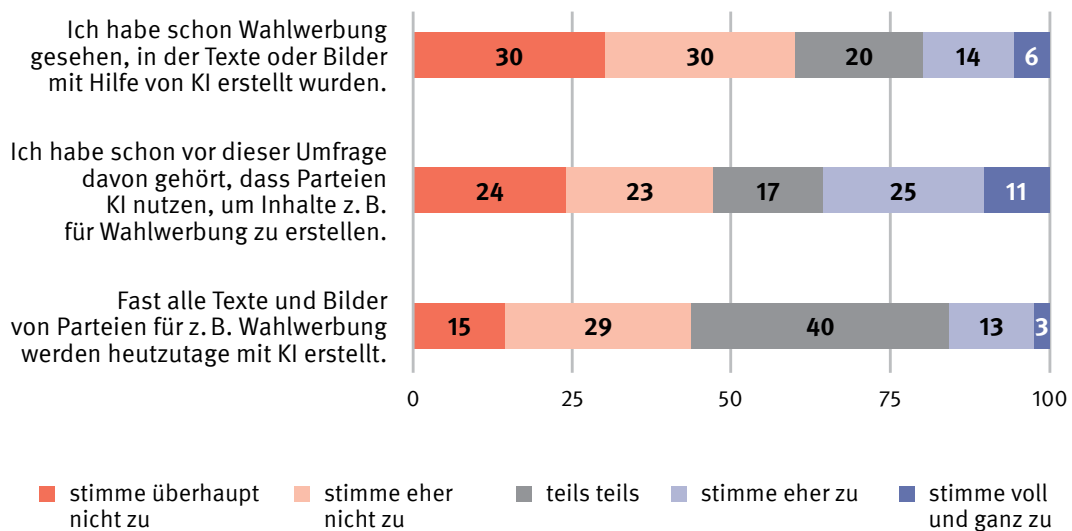
der politischen Kampagnenarbeit von Parteien hin, jedoch besteht bei vielen Menschen auch Unklarheit über den konkreten Einsatz solcher Technologien.

3.3.2 Allgemeine Beurteilung des KI-Einsatzes

Die Vorstellung vom KI-Einsatz in politischen Kampagnen löst bei den Befragten eher Skepsis aus. In keinem Fall überwiegen positive Einschätzungen (siehe Abbildung 14). So finden mehr als die Hälfte der Befragten, dass der Einsatz von KI in Kampagnen (eher) gefährlich (59 %), (eher) nicht wünschenswert (56 %) und (eher) unangebracht ist (51 %). Darüber hinaus

Abbildung 13:

Kenntnisse von KI in der Politik (in %)



Fragetext: „Wir werden Ihnen nun Fragen zu Ihren bisherigen Erfahrungen mit KI in der politischen Kommunikation stellen. Wie sehr stimmen Sie folgenden Aussagen zu?“

Quelle: Eigene Darstellung.

empfinden knapp die Hälfte der Befragten die politische KI-Nutzung als (eher) inakzeptabel (46 %). Die Nutzung von KI in politischen Kampagnen löst bei 41 Prozent Angst aus; ebenso viele finden den Einsatz von KI (eher) unfair oder sind unentschlossen.

Andererseits löst der Einsatz von KI bei einigen Befragten weniger Sorgen als positive Assoziationen aus. So empfinden immerhin über ein Drittel (37 %) KI als nützlich für Kampagnen.

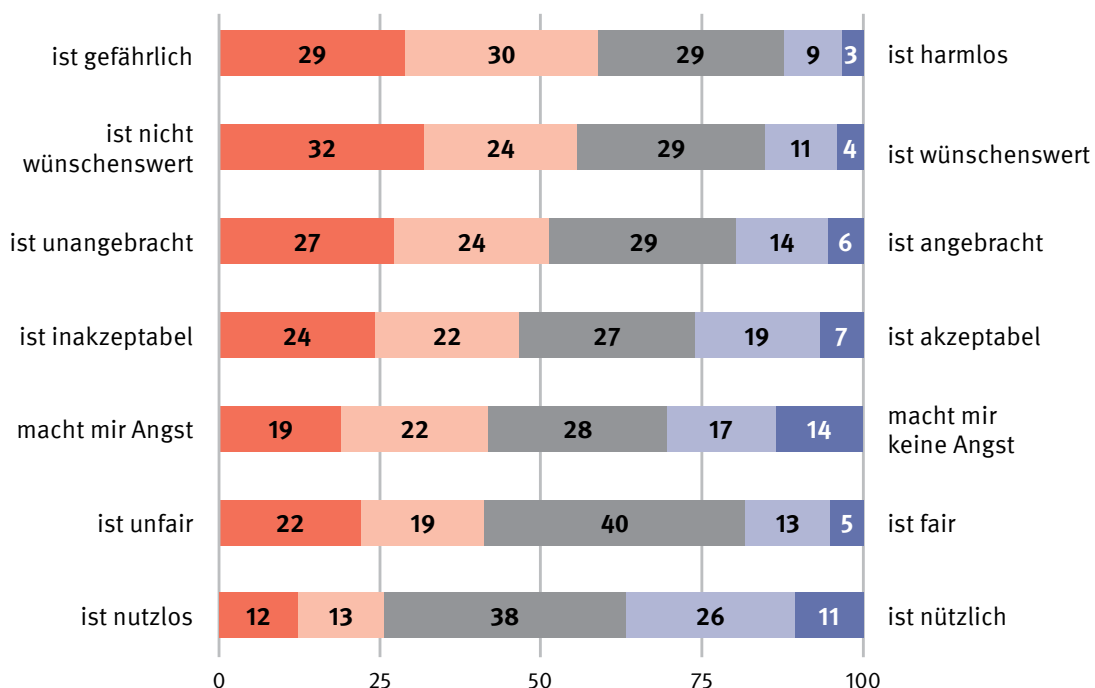
Ebenso viele sind diesbezüglich ambivalent. Darüber hinaus hat knapp ein Drittel der Befragten keine Angst vor einem KI-Einsatz in Politik-kampagnen (31%) und findet diesen sogar akzeptabel (26 %).

3.3.3 Bewertung von Parteien, die KI einsetzen

Viele Deutsche begegnen dem Einsatz von KI durch politische Parteien mit Misstrauen (siehe Abbildung 15). Mehr als drei Viertel der Be-

Abbildung 14:
Bewertung von KI in der Politik (in %)

Frage-*text*: „Wie beurteilen Sie, ganz persönlich, den Einsatz von KI in politischen Kampagnen (z. B. im Wahlkampf)? Der Einsatz von KI zur Erstellung von politischen Kampagnen ...“



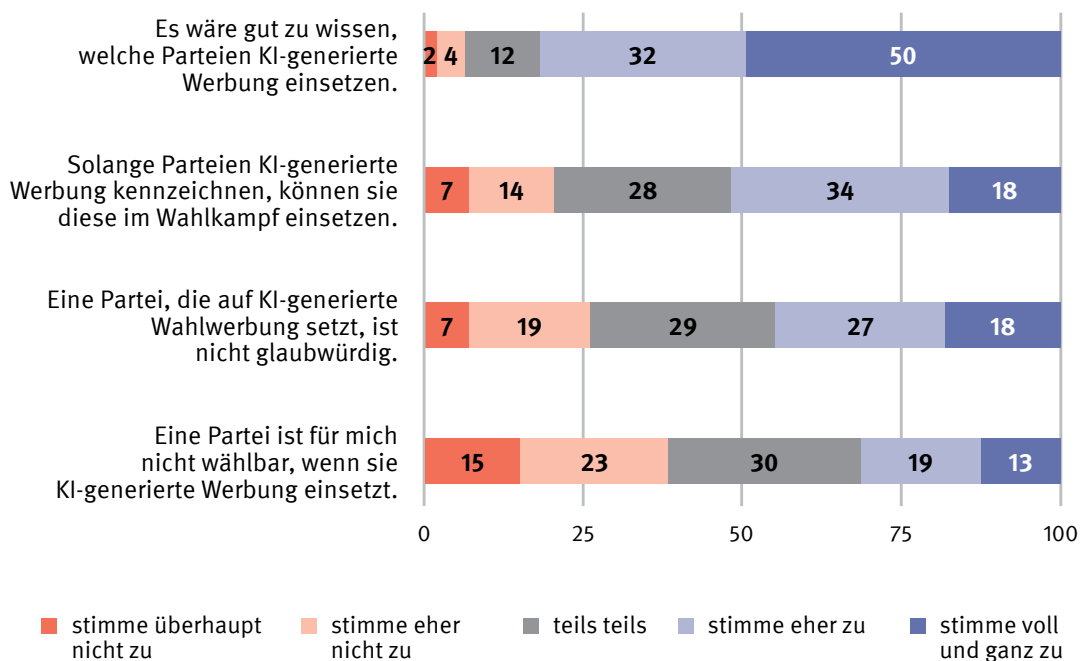
Quelle: Eigene Darstellung.

fragten (82%) gab an, wissen zu wollen, welche Parteien KI einsetzen. Fast die Hälfte der Befragten (45%) sieht beim Einsatz von KI die Glaubwürdigkeit einer Partei gefährdet und rund ein Drittel hält eine Partei dann für nicht mehr wählbar (32%). Es gibt aber auch andere Stimmen. So findet die Hälfte der Befragten, dass der Einsatz von KI in der Wahlwerbung akzeptabel ist, sofern dies von den Parteien transparent gekennzeichnet wird (52%). Diese Einschätzungen verdeutlichen eine differenzierte, aber überwiegend kritische Haltung der deutschen Bevölkerung gegenüber Parteien, die KI einsetzen.

3.3.4 Akzeptanz bei unterschiedlichen Aufgaben

Die deutsche Bevölkerung bewertet den Einsatz von KI in der Kommunikation politischer Akteur*innen je nach Aufgabe sehr unterschiedlich (siehe Abbildung 16). Eine breite Mehrheit stimmt dem Einsatz von KI zur Kontrolle von Rechtschreibung und Grammatik in Texten (80%) sowie zur sprachlichen Überarbeitung (65%) (eher) zu. Auch die Zusammenfassung und Übersetzung komplexer Texte in einfache Sprache wird von über der Hälfte (57%) positiv gesehen. Allerdings sinkt die Akzeptanz, wenn es um Aufgaben geht, die den direkten Kontakt mit Bürger*innen be-

Abbildung 15:
Einschätzung von Parteien, die KI einsetzen (in %)



Fragestext: „Angenommen im nächsten Bundestagswahlkampf würden Parteien KI-generierte Werbung verwenden. Inwiefern stimmen Sie diesen Aussagen zu?“

Quelle: Eigene Darstellung.

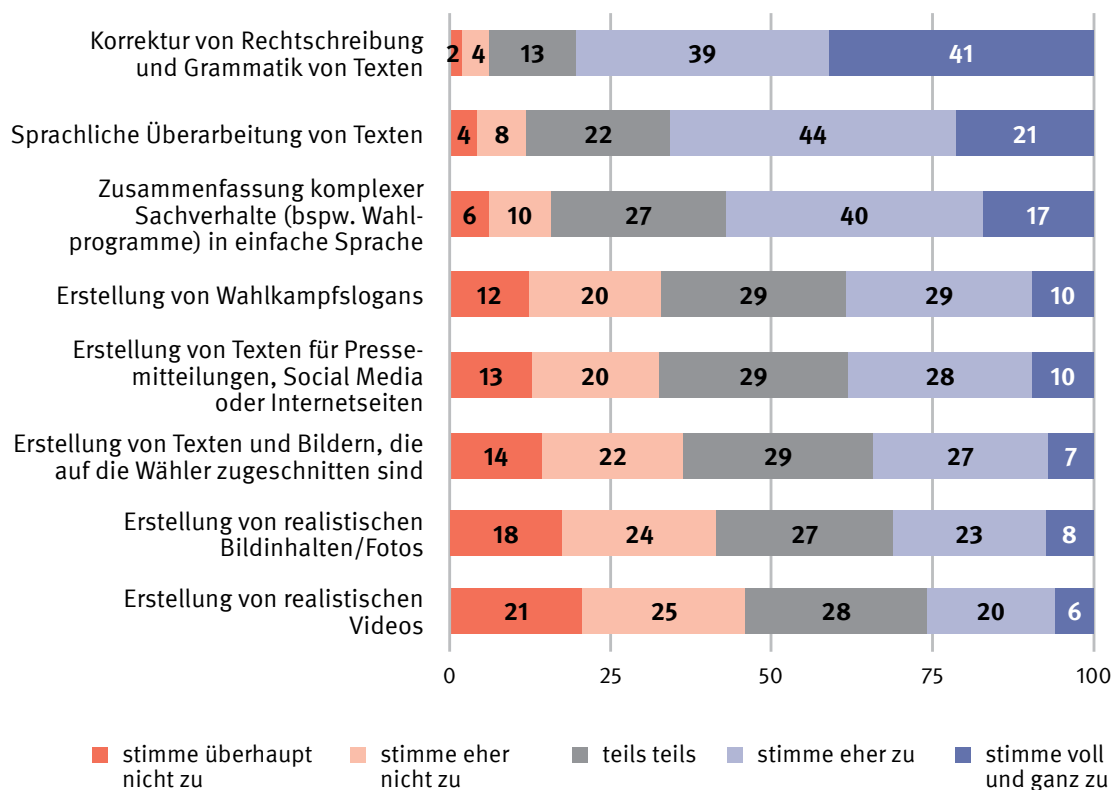
treffen. So sinkt die Zustimmung zum Einsatz von KI zur Erstellung von Wahlkampflogos (39 %) oder von Pressemitteilungen, Social-Media-Content oder Internetseiten (38 %). Skepsis ruft der Einsatz von KI auch bei der Erstellung von Texten und Inhalten hervor, die auf die Wähler*innen zugeschnitten sind; rund 36 Prozent stimmen einem solchen Einsatz (eher) nicht zu. Am wenigsten akzeptiert wird der Einsatz von KI für die Erstellung von realistischen Bildern (42 %) und Videos (46 %). Jeweils fast die Hälfte der Befragten findet, dass diese Anwendung von KI eher oder überhaupt nicht akzeptabel ist.

3.3.5 Chancen und Risiken

Nur wenige Befragte verbinden mit dem KI-Einsatz in politischen Kampagnen Chancen (siehe Abbildung 17). Zwar erkennt etwa die Hälfte das Potenzial von KI, Wahlwerbung gezielter auf die Interessen der Wähler*innen zuzuschneiden (47 %). Jedoch besteht Unklarheit, ob KI in der Lage ist, die Qualität politischer Diskussionen zu verbessern oder sachlichere Debatten in Wahlkämpfen zu fördern. Obwohl rund 35 Prozent der Befragten sagen, dass der Wahlkampf durch KI weniger emotional wird, zweifeln daran fast genauso viele (36 %) oder sind unentschieden (30 %).

Abbildung 16:

Akzeptanz von KI in der Politik nach Einsatzgebieten (in %)



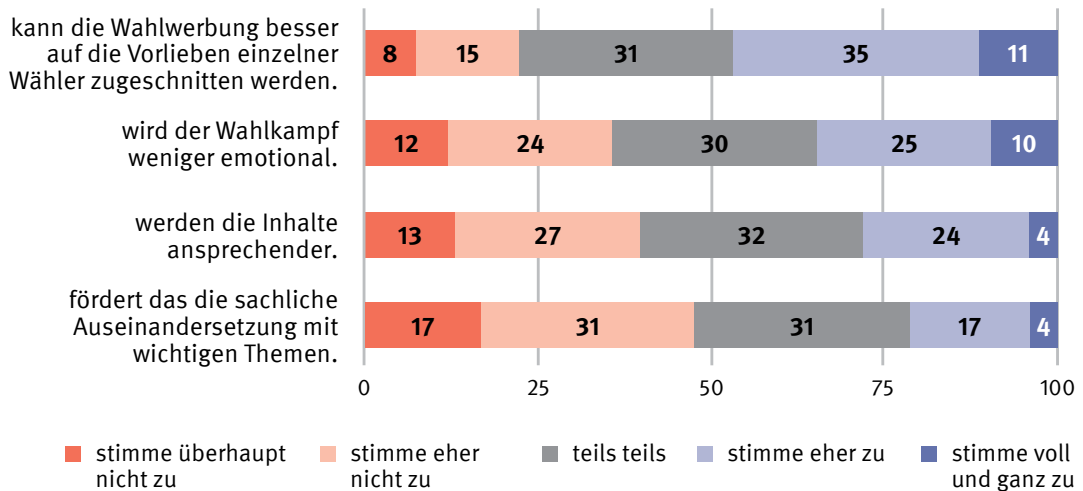
*Frage*text: „Nun geht es um die Möglichkeiten des Einsatzes von KI in der Kommunikation von Politikern und Parteien. Bitte geben Sie an, für wie akzeptabel Sie den Einsatz von KI für bestimmte Aufgaben halten.“

Quelle: Eigene Darstellung.

Abbildung 17:

Chancen eines KI-Einsatzes in der Politik (in %)

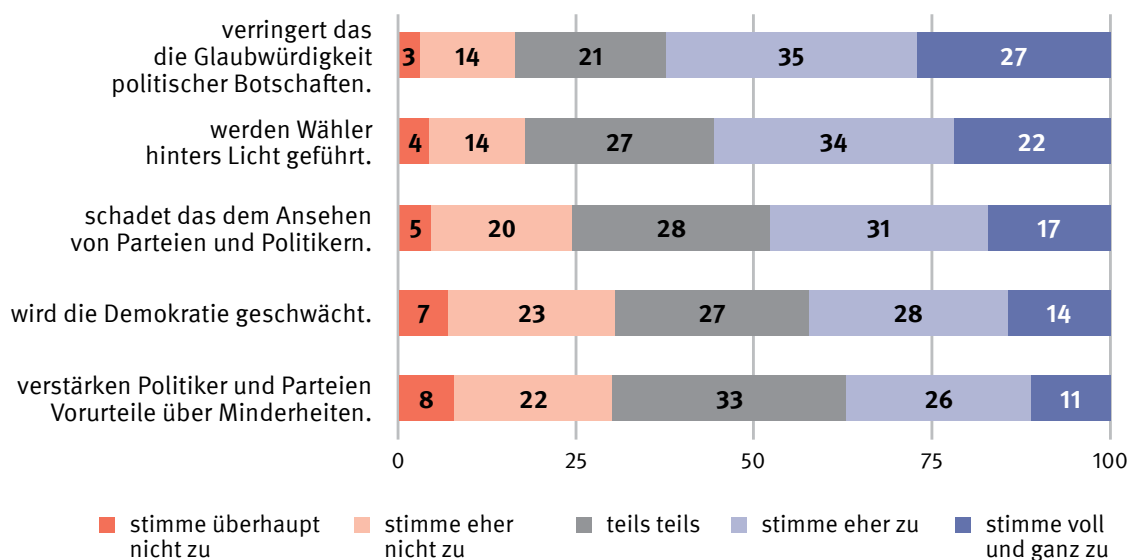
*Frage*text: „Wir werden Ihnen nun einige Fragen zu Ihrer persönlichen Einschätzung zum Einsatz von KI in politischen Kampagnen stellen. Bitte geben Sie an, wie sehr Sie den folgenden Aussagen zustimmen. Wenn Parteien und Politiker in Zukunft immer öfter KI für das Erstellen von politischer Werbung verwenden, ...“



Quelle: Eigene Darstellung.

Abbildung 18:

Risiken eines KI-Einsatzes in der Politik (in %)



*Frage*text: „Wir werden Ihnen nun einige Fragen zu Ihrer persönlichen Einschätzung zum Einsatz von KI in politischen Kampagnen stellen. Bitte geben Sie an, wie sehr Sie den folgenden Aussagen zustimmen. Wenn Parteien und Politiker in Zukunft immer öfter KI für das Erstellen von politischer Werbung verwenden, ...“

Quelle: Eigene Darstellung.

Während etwas mehr als ein Viertel eine Chance in KI sieht, die Botschaftsinhalte ansprechender für die Wähler*innen zu gestalten, widerspricht dem mehr als ein Drittel (40 %). Diese Skepsis begründet sich auch in der Befürchtung, dass der Einsatz von KI nicht zu einer sachlicheren Auseinandersetzung mit politischen Themen führt (48 %).

Statt mehrheitlich Hoffnung auf neue Möglichkeiten zu wecken, löst der Einsatz von KI in der politischen Kampagnenarbeit also eher gemischte Gefühle aus. Viele Befragte zeigten sich skeptisch und sehen im Einsatz gar konkrete Gefahren für politische Akteur*innen und die Demokratie als Ganzes (siehe Abbildung 18). Knapp zwei Drittel der Befragten sind (eher) überzeugt, dass KI die Glaubwürdigkeit politischer Botschaften verringert (62 %), wobei nur 17 Prozent diese Befürchtung (eher) nicht teilen. Besonders auffällig ist, dass etwas mehr als die Hälfte der Befragten (56 %) der Meinung ist, Wähler*innen könnten durch den Einsatz von KI in die Irre geführt werden, was ein erhöhtes Bewusstsein für Manipulationsrisiken durch KI in Kampagnen andeutet. Darüber hinaus teilt knapp die Hälfte der Befragten die Einschätzung, dass KI das Ansehen von Parteien und Politiker*innen gefährdet (48 %). Nur wenige sehen diese drei Risiken (eher) nicht. Ob die Demokratie durch KI in Gefahr gerät, beurteilen die Befragten unterschiedlich. Während rund 42 Prozent KI als Gefahr für Demokratien halten, widersprechen dem drei von zehn Befragten (30 %). Mehr als ein Drittel der Befragten (37 %) ist überzeugt, dass der Einsatz von KI Vorurteile gegenüber Minder-

heiten verstärken könnte, wobei dem jeweils fast genauso viele widersprechen und/oder unentschieden sind. Auch wenn das Meinungsbild teilweise variiert, ist doch zu sehen, dass die deutsche Bevölkerung dem Einsatz von KI für politische Kampagnen überwiegend kritisch gegenübersteht. Insbesondere in Bezug auf die Integrität und Fairness des demokratischen Prozesses.

3.3.6 Bewertung von Regulierungsvorschlägen

Breite Unterstützung gibt es für Vorschläge zur Regulierung des Einsatzes von KI in politischen Kampagnen (siehe Abbildung 19). Eine große Mehrheit der Befragten (83 %) ist der Meinung, dass alle mit KI erstellten politischen Inhalte gekennzeichnet werden sollten. Dies gilt besonders für realistische Bilder von Personen oder Ereignissen, bei denen fast 85 Prozent der Befragten dieser Aussage voll und ganz oder eher zustimmen. Auch für eindeutig unrealistische oder künstlich wirkende Bilder fordern fast ebenso viele eine klare Kennzeichnung (78 %). Darüber hinaus stimmen drei Viertel der Befragten (eher) zu, dass KI-generierte Inhalte vor ihrer Veröffentlichung von unabhängigen Expert*innen überprüft werden sollten, um sicherzustellen, dass sie der Wahrheit entsprechen. Ein freiwilliger Verzicht oder ein striktes Verbot findet hingegen keine eindeutige Mehrheit. So fordert zwar knapp die Hälfte der Deutschen (48 %), Parteien und Politiker*innen sollten freiwillig auf den Einsatz von KI verzichten. Hinsichtlich eines kompletten Verbots von KI in politischen Kampagnen zeigen sich die Befragten jedoch

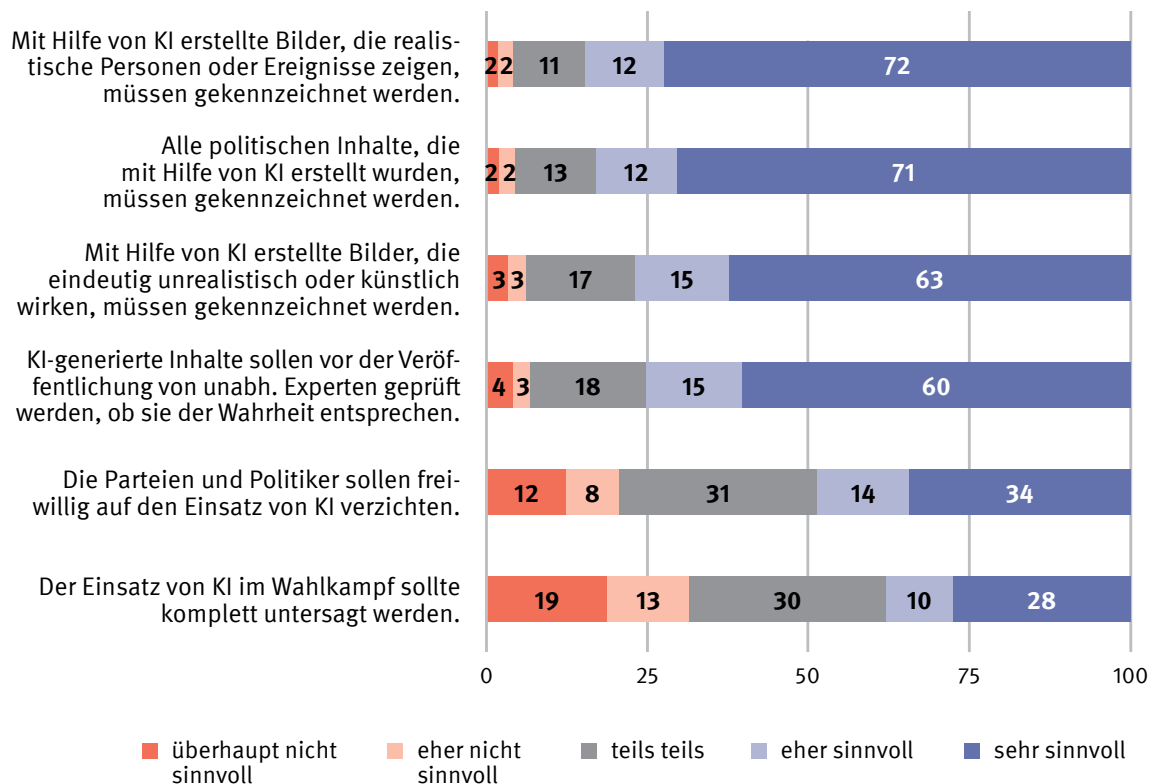
ambivalent: knapp vier von zehn Befragten (38 %) stimmen diesem zumindest eher zu. Es bestehen somit grundlegende Bedenken gegenüber der Nutzung von KI in der politischen Kommunikation, die bei vielen Befragten mit dem Wunsch nach strengeren Regulierungen oder sogar einem vollständigen Verbot einhergeht. Gleichzeitig spricht sich aber auch eine große Anzahl der Befragten gegen diese beiden strikten Regulierungsmöglichkeiten aus (Verzicht: 20 %; Verbot: 31 %).

3.4 Zwischenfazit: Wenig Berührungspunkte mit KI-Kampagnen, aber weit verbreitete Skepsis und Sorgen

Die Forschungsergebnisse der repräsentativen Befragung zur Wahrnehmung und Akzeptanz der KI-Nutzung in der politischen Kommunikation sowie zum wahrgenommenen Regulierungsbedarf zeichnen ein differenziertes Bild. Während die Befragten zwar durchaus Potenziale erkennen, überwiegen Skepsis und Bedenken hin-

Abbildung 19:

Regulierung des Einsatzes von KI in der Politik (in %)



*Frage*text: „Zum Abschluss geht es nun um Maßnahmen und Regelungen für den Einsatz von KI in den kommenden Wahlkämpfen. Wie sinnvoll finden Sie die folgenden Vorschläge?“

Quelle: Eigene Darstellung. Abweichungen von 100 basieren auf Rundungen.

sichtlich möglicher negativer Auswirkungen auf die Integrität und Qualität des politischen Diskurses. Die Befürwortung regulatorischer Maßnahmen, insbesondere in Bezug auf Transparenz und unabhängige Überprüfung, unterstreicht die Besorgnis über gesteigerte Manipulationen und Fehlinformationen als potenzielle Gefahren. Gleichzeitig zeigt die differenzierte Bewertung verschiedener Einsatzbereiche von KI, dass eine pauschale Ablehnung oder Befürwortung zu kurz greift.

Die Ergebnisse der Befragung lassen sich wie folgt entlang der aufgeworfenen Forschungsfragen zusammenfassen.

Fragestellung 1: Welche Kenntnisse haben Bürger*innen über KI in politischen Kampagnen?

- **Geringe direkte Berührungspunkte mit KI:** Eine Mehrheit der Befragten gibt an, noch keine KI-generierten Inhalte in der politischen Wahlwerbung gesehen zu haben (z.B. Texte oder Bilder).
- **Unsicherheit über KI-Einsatz durch Parteien:** Knapp die Hälfte der Befragten ist ungewiss, ob politische Parteien überhaupt KI für die Erstellung von Wahlwerbung nutzen.
- **Wahrnehmung begrenzter KI-Nutzung:** Nur wenige Befragte vermuten, dass Parteien nahezu alle Inhalte, wie Texte oder Bilder, mit KI erstellen.
- **Diskrepanz zwischen Nutzung und Wahrnehmung:** Es gibt eine Lücke zwischen der tatsächlichen Nutzung von KI durch politische Parteien und der Wahrnehmung dieser Nutzung in der Öffentlichkeit.

Fragestellung 2: Wie beurteilen Bürger*innen den Einsatz generativer KI in politischen Kampagnen?

- **Vorherrschend negative Einstellungen:** Über die Hälfte der Befragten empfindet den Einsatz von KI in politischen Kampagnen als gefährlich, nicht wünschenswert oder unangebracht. Weniger als die Hälfte sieht den KI-Einsatz als inakzeptabel oder unfair an.
- **Aber auch positive Haltungen und Gelassenheit:** Ein wesentlicher Teil der Befragten erkennt die Nützlichkeit von KI für politische Kampagnen. Mehr als ein Drittel äußert keine Angst vor dem KI-Einsatz und bewertet ihn sogar als akzeptabel.
- **Ambivalente Meinungen:** Die Meinungen der Befragten sind gespalten und reichen von Ablehnung bis hin zu Akzeptanz. Dies verdeutlicht die kontroversen Haltungen in der öffentlichen Debatte über den Einsatz von KI in der politischen Kommunikation.

Fragestellung 3: Wie bewerten Bürger*innen Parteien, die generative KI einsetzen?

- **Starkes Bedürfnis nach Transparenz:** Mehr als drei Viertel der Befragten möchten wissen, wenn Parteien KI in ihren Kampagnen einsetzen. Ein transparenter Umgang mit KI durch politische Akteur*innen ist ihnen daher wichtig.
- **Glaubwürdigkeitsverlust durch KI-Einsatz:** Fast die Hälfte der Befragten sieht die Glaubwürdigkeit von Parteien gefährdet, wenn KI in politischen Kampagnen verwendet wird. Mehr als ein Viertel würde eine solche Partei nicht mehr wählen.

- **Potenziale zur Akzeptanzsteigerung:** Etwa die Hälfte der Befragten bewertet den KI-Einsatz in der Wahlwerbung als akzeptabel, wenn dieser von den Parteien klar gekennzeichnet wird.

Fragestellung 4: Inwieweit akzeptieren Bürger*innen den Einsatz generativer KI für bestimmte Aufgaben bei der Wählerkommunikation?

- **Hohe Akzeptanz für technische und unterstützende Aufgaben:** Die Mehrheit der Befragten akzeptiert den KI-Einsatz für Aufgaben wie Rechtschreib- und Grammatikprüfung, sprachliche Überarbeitung von Texten und die Vereinfachung komplexer Inhalte.
- **Sinkende Akzeptanz bei inhaltlicher Gestaltung:** Skepsis herrscht gegenüber dem Einsatz von KI für kreative und strategische Aufgaben wie die Erstellung von Wahlkampfslogans, Pressemitteilungen, Social-Media-Inhalten oder Internetseiten.
- **Kritik an KI-generierten personalisierten Inhalten:** Besonders kritisch wird der Einsatz von KI für personalisierte Inhalte und Wählerkommunikation gesehen.
- **Ablehnung von KI-generierten visuellen Inhalten:** Fast die Hälfte der Befragten lehnt die Erstellung realistischer Bilder und Videos durch KI ab.
- **Verstärkte Bedenken bei komplexeren Anwendungen:** Je stärker KI in die inhaltliche Gestaltung und Personalisierung politischer Botschaften eingreift, desto größer sind die Vorbehalte.

Fragestellung 5: Welche Chancen und Risiken sehen Bürger*innen beim Einsatz generativer KI in politischen Kampagnen?

- **Überwiegende Skepsis gegenüber KI in Kampagnen:** Die deutsche Bevölkerung bewertet den Einsatz von KI in politischen Kampagnen eher kritisch und erkennt mehr Risiken als Chancen.
- **Wahrnehmung erheblicher Risiken:** Eine Mehrheit der Befragten sieht durch den KI-Einsatz in Kampagnen die Glaubwürdigkeit politischer Botschaften gefährdet und befürchtet, dass Wähler*innen in die Irre geführt werden.
- **Anerkennung bestimmter Chancen:** Allerdings sieht etwa die Hälfte der Befragten auch das Potenzial, eine gezieltere Ansprache von Wähler*inneninteressen durch den Einsatz von KI zu erreichen.
- **Geringe Zustimmung zu weiteren Vorteilen:** Nur eine Minderheit der Befragten glaubt, dass KI die Qualität politischer Diskussionen verbessert, zu sachlicheren Debatten beiträgt oder dabei hilft, Inhalte ansprechender zu gestalten.
- **Bedrohung demokratischer Werte:** Knapp die Hälfte der Befragten bewertet den KI-Einsatz als eine Gefahr für das Ansehen von Parteien, Politiker*innen und sogar für die Demokratie insgesamt.

Fragestellung 6: Wie bewerten Bürger*innen verschiedene Formen der Regulierung von generativer KI in politischen Kampagnen?

- **Breite Unterstützung für Transparenzmaßnahmen:** Eine überwältigende Mehrheit der

Befragten fordert eine klare Kennzeichnung aller mit KI erstellten politischen Inhalte, insbesondere bei realistischen Bildern von Personen oder Ereignissen. Auch für künstlich wirkende Bilder wird eine Kennzeichnung befürwortet, wenn auch in geringerer Intensität.

- **Prüfung durch unabhängige Expert*innen:** Rund drei Viertel der Befragten sprechen sich dafür aus, dass KI-generierte Inhalte vor ihrer Veröffentlichung durch unabhängige Expert*innen geprüft werden sollten.
- **Rückhalt für strengere Restriktionen:** Ein erheblicher Teil der Befragten unterstützt sogar härtere Maßnahmen, wie einen freiwilli-

gen Verzicht oder ein komplettes Verbot des KI-Einsatzes in politischen Kampagnen.

- **Hohes Bedürfnis nach Schutz:** Die breite Zustimmung zu Transparenz und Restriktionen zeigt, dass viele Deutsche die Risiken des KI-Einsatzes in politischen Kampagnen als schwerwiegend einstufen und Maßnahmen zum Schutz demokratischer Werte und Verfahren wünschen.
- **Konsistenz mit regulatorischen Rahmenbedingungen:** Die Forderungen der Befragten stehen im Einklang mit den Regulierungsansätzen des europäischen AI Acts und den Selbstverpflichtungen einiger Parteien und Plattformen.

4 Online-Experiment zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern

Nachdem wir in Kapitel 3 gegenwärtige Wissensstände und Einstellungen der deutschen Bevölkerung zum Einsatz generativer KI in politischen Kampagnen dargestellt haben, wechseln wir in unserer zweiten empirischen Untersuchung die Perspektive und das Untersuchungsdesign. Mit einem Online-Experiment untersuchen wir die Wahrnehmung und Wirkung KI-generierter Bilder in politischen Kampagnen. Dabei möchten wir Erkenntnisse darüber gewinnen, inwiefern KI-generierte Kampagnenbilder überhaupt als solche erkannt werden, wie diese bewertet werden und welche Reaktionen sie auslösen. Zudem betrachten wir, wie sich die Rezeption von KI-Bildern auf die Einstellungen der Rezipient*innen zum Einsatz von KI in politischen Kampagnen auswirkt. Dabei untersuchen wir, inwiefern die Bewertung von Parteien und die generelle Wahrnehmung von KI als Risiko für die Demokratie beeinflusst wird.

Analog zu Kapitel 3 diskutieren wir im Folgenden zunächst ausführlicher die Forschungsfragen zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern, denen wir in diesem Kapitel nachgehen und ordnen diese in den aktuellen Stand der Forschung ein (4.1). Anschließend stellen wir das Design unserer Untersuchung vor (4.2), diskutieren darauf aufbauend die Ergebnisse des Experiments (4.3) und fassen die zentralen Erkenntnisse in einem Zwischenfazit zusammen (4.4).

4.1 Erkenntnisinteresse und Forschungsfragen

Im folgenden Abschnitt diskutieren wir ausführlich das Untersuchungsziel und die übergeordneten Fragestellungen unseres Experiments. Zunächst interessiert uns die unmittelbare Wahrnehmung von KI-generierten Kampagnenbildern. In einem zweiten Schritt möchten wir herausfinden, wie die Rezeption KI-generierter Bilder die grundlegenden Einstellungen der Menschen zum KI-Einsatz beeinflusst. Dabei untersuchen wir, inwieweit die Beurteilung der Inhalte und die mittelfristigen Einstellungen der Rezipient*innen von den Absender*innen der KI-Botschaften, dem Botschaftsinhalt sowie der Kennzeichnung durch Transparenzhinweisen beeinflusst wird. Die grundlegenden Konzepte und theoretischen Annahmen unserer Forschungsfragen betten wir in den aktuellen Forschungsstand zur Wahrnehmung und Wirkung KI-generierter Bilder ein.

Fragestellung 7: Welchen Einfluss haben Bildinhalte und Transparenzhinweise auf die Wahrnehmung KI-generierter Kampagnenbilder?

Unsere Befragung hat gezeigt, dass viele Menschen laut eigener Aussagen bisher wenig Berührungspunkte mit KI-generierten Inhalten hatten (siehe Kapitel 3.3.1). Im Lichte dieser Einschätzung ist es relevant herauszufinden, ob Menschen überhaupt KI-generierte Kampagnen-

bilder erkennen und von echten unterscheiden können. Darüber hinaus sind Antworten auf diese Frage relevant, weil die Erkennung von KI-generierten Bildern laut Studien einen Einfluss auf die durch die Botschaft ausgelösten Emotionen und Bewertungen hat und damit auch Einstellungen prägen können (sogenannte Downstream-Effekte; Appel & Prietzel, 2022).

Die bisherige Forschung zur Wahrnehmung von KI-generierten Inhalten hat sich größtenteils mit Inhalten beschäftigt, die durch ältere KI-Modelle erstellt wurden. Die Qualität dieser älteren Modelle war teils deutlich geringer und die Inhalte konnten daher leichter als künstlich generiert entlarvt werden. Die jüngsten Weiterentwicklungen dieser Modelle haben die Qualität der generierten Inhalte deutlich verbessert und erlauben die Erzeugung hochrealistischer Kampagnenbilder. Inwieweit Menschen diese sehr realitätsnahen KI-Inhalte erkennen und wie sie diese bewerten ist bis dato noch nicht erforscht. Daher wollen wir mit unserem Experiment erste Erkenntnisse darüber gewinnen. Der allgemeinen Frage nach dem Einfluss von Bildinhalten und Transparenzhinweisen auf die Wahrnehmung KI-generierter Kampagnenbilder gehen wir dabei mit Hilfe von drei kleinteiligen Forschungsfragen nach.

Fragestellung 7.1: Werden KI-generierte Kampagnenbilder als solche erkannt?

Mit der zunehmenden Verbreitung und Verfeinerung von KI-Technologien wird es immer schwieriger, zwischen authentischen und künstlich erzeugten Inhalten zu unterscheiden (u. a. Frank et al., 2023; Groh et al., 2024; Köbis et al., 2021).

Während das menschliche visuelle System evolutionär darauf ausgelegt ist, natürliche Szenen und Objekte zu erkennen und zu interpretieren, hat es keine spezifischen Mechanismen entwickelt, um künstlich erzeugte Bilder zu identifizieren, die bewusst darauf ausgelegt sind, natürlich zu erscheinen (siehe „Realismus-Heuristik“ von Sundar, 2008). Beim Rezeptionsprozess können wichtige Hinweise auf den KI-Ursprung eines Bildes leicht übersehen werden, insbesondere wenn diese subtil sind (beispielsweise bei der Bildgenerierung entstehende Verzerrungen, Ungereimtheiten oder Fehler im Bild wie sechs Finger an den Händen von Menschen; siehe Infobox 6) oder sich in Bereichen befinden, die nicht im Fokus unserer Aufmerksamkeit liegen (beispielsweise KI-Kennzeichnungen).

Zur spezifischen Frage der Erkennung von KI-generierten politischen Kampagnenbildern liegen bisher keine Studien vor. Jedoch offenbaren Erkenntnisse zu KI-Texten und unpolitischen KI-Bildern oder politischen Deepfake-Videos große Schwierigkeiten bei der Unterscheidung zwischen „echt“ und „unecht“: Im Bereich der Texterkennung fanden Kreps et al. (2022), dass die meisten Teilnehmenden nicht in der Lage waren, zwischen KI-generierten und menschlich geschriebenen Nachrichtenartikel zu unterscheiden. Frank et al. (2023) führten eine länderübergreifende Untersuchung in den USA, Deutschland und China durch und kommen zu dem Schluss, dass Menschen KI-generierte Texte, Bilder und Audios nur unzureichend erkennen. Insbesondere bei Bildern war die Erkennungsrate am schlechtesten. Darüber hinaus zeigen Nightingale und Farid (2022), dass KI-generierte

Gesichter für Menschen kaum von echten Gesichtern unterscheidbar sind. Weitere Studien zur Erkennung von politischen Deepfake-Videos bestätigen diese Ergebnisse größtenteils und stellen fest, dass nur wenige Menschen in der Lage sind, Deepfake-Videos als solche zu erkennen, selbst wenn die KI-generierten Inhalte technische oder anatomische Fehler aufweisen (Dobber et al., 2021; Hameleers et al., 2024; Köbis et al., 2021; Lovato et al., 2024).

Folgende drei Faktoren untersuchen wir, die die Erkennung von KI-generierten Inhalten beeinflussen können:

(1) Zunächst geht es uns um Botschaftsinhalte wie *eine positive oder negative Bildsprache* durch visuelle Darstellungen oder Texte (sogenannte Valenz; Leonhard & Bartsch, 2020; Scheimer, 2010). Welche Rolle diese für die Erkennung von KI-generierten Bildern spielen, ist weitestgehend unbekannt. Nur Nightingale und Farid (2022) zeigen, dass lächelnde positive Gesichter häufiger als echt eingestuft werden, was darauf hindeutet, dass positive Bildinhalte die Erkennung von KI-generierten Inhalten erschweren könnten.

(2) Außerdem interessieren uns Hinweise auf den KI-Ursprung. *Aufklärung(stexte)* über die Möglichkeiten von KI zur Erstellung realitätsnaher Bilder zeigten in einigen Studien positive Effekte auf die Erkennungsfähigkeit. Hwang et al. (2021) beobachteten, dass Bildungsvideos über Deepfakes die Fähigkeit der Teilnehmenden, KI-generierte Inhalte zu identifizieren,

verbesserten. Auch ein Trainingsprogramm zur Detektion von KI-generierten Inhalten erhöht die Erkennungsrate deutlich (Tahir et al., 2021). Dagegen stellten Köbis et al. (2021) oder Nightingale und Farid (2022) fest, dass selbst mit *Vorwissen* über KI-Manipulationen die Unterscheidung zwischen echten und synthetischen Bildern für viele Menschen eine Herausforderung blieb. Dieses Ergebnis zeigt sich auch in qualitativen Fokusgruppen, in denen Menschen die Echtheit von Bildern oft nicht in Frage stellten, obwohl sie durch Diskussionen auf Möglichkeiten der Manipulation aufmerksam gemacht wurden (Kasra et al., 2018).

Für *KI-Disclaimer* weisen beispielsweise Lu und Yuan (2024) nach, dass sie das Verständnis des Publikums und seine Fähigkeit, ein Deepfake zu erkennen, positiv beeinflussen. Jedoch sind die Forschungsergebnisse für KI-Kennzeichnungen oder Hinweislables nicht eindeutig. So kamen Appel und Prietzel (2022) zu dem Ergebnis, dass Versuchsteilnehmende trotz Hinweisen auf mögliche Manipulationen Schwierigkeiten hatten, Deepfakes zuverlässig zu identifizieren.

(3) Letztlich sind wir auch am Einfluss von individuellen Eigenschaften der Menschen interessiert. So identifizierten Frank et al. (2023) allgemeines *Vertrauen*, *Medienkompetenz* sowie *KI-Kenntnisse* als signifikante Einflussfaktoren für die KI-Erkennung. Kreps et al. (2022) stellten fest, dass Teilnehmende KI-generierte Inhalte eher als solche identifizieren, wenn diese nicht mit ihren politischen Ansichten übereinstimmen (siehe auch Shen et al., 2019).

Darauf aufbauend untersuchen wir für möglichst realitätsnahe KI-Bilder im politischen Kampagnenkontext, wie gut Menschen in der Lage sind, diese Bilder als künstlich zu identifizieren und von echten Aufnahmen zu unterscheiden.

Welche Emotionen (Fragestellung 7.2) und Bewertungen (Fragestellung 7.3) lösen KI-generierte Kampagnenbilder aus?

Emotionen sind ein wesentlicher Bestandteil von Kampagnenkommunikation und können bei Wähler*innen insbesondere durch visuelle Botschaften hervorgerufen werden (Geise, 2011; Maurer, 2016). Im Vergleich zu Texten haben Bilder in der politischen Kommunikation den Vorteil, dass sie schneller verarbeitet werden und direkt emotionale Reaktionen auslösen können (Geise, 2011; Leonhard & Bartsch, 2020; Maurer, 2016). Im Sinne der sogenannten „emotional contagion“ (Masch & Gabriel, 2020) können Emotionen wie Hoffnung, Angst oder Wut bei den Betrachter*innen erzeugt werden, etwa durch die Verwendung bestimmter Farben, Symbole oder Gesichtsausdrücke. Das Ziel solcher emotionalisierenden Strategien ist es aber nicht nur, kurzfristige Emotionen und Aufmerksamkeit bei den Wähler*innen auszulösen, sondern auch ihre Einstellungen und Verhaltensweisen zu beeinflussen (Grüning & Schubert, 2022; Masch & Gabriel, 2020; Maurer, 2014).

Eine wichtige Rolle spielt dabei die Bewertung der Botschaften. Dabei sind die wahrgenommene Natürlichkeit, Überzeugungskraft und Unaufdringlichkeit von Botschaften wichtige Faktoren, die Akzeptanz und Wirkung von Kampagnen-

botschaften steigern oder reduzieren können (Friestad & Wright, 1994; Maurer, 2014). Natürlich wirkende Botschaften werden seltener als manipulativ wahrgenommen, was das Vertrauen in die Absenderpartei stärkt. Gleichzeitig erhöhen überzeugende Botschaften, die durch starke Argumente gestützt sind, die Wahrscheinlichkeit, dass Bürger*innen ihre Einstellungen ändern oder festigen (Petty & Cacioppo, 1986). Unaufdringliche Botschaften rufen zudem weniger Reaktanz hervor, da sie den Rezipient*innen nicht das Gefühl vermitteln, manipuliert zu werden (Brehm & Brehm, 1981; Friestad & Wright, 1994).

Mit den neuen technologischen Entwicklungen generativer KI steigern sich die Möglichkeiten, emotionale Wirkungen und gezielte Bewertungen durch Kampagnenbilder hervorzurufen, noch einmal beträchtlich. Mit den richtigen Prompts können Bildinhalte wie etwa positive oder negative Gesichtsausdrücke, Farbgebung oder Stimmungen exakt an die gewünschten emotionalen Reaktionen der Zielgruppe angepasst werden und durch den hohen Realitätsgrad als natürlich, überzeugend und unaufdringlich wahrgenommen werden (siehe Kapitel 2.1.2; Infobox 2). Wir untersuchen die folgenden vier Faktoren, die einen Einfluss auf die emotionalen Reaktionen und Bewertungen haben können:

(1) Untersuchungen zeigen, dass *KI-generierte Bilder*, insbesondere solche mit Gesichtern, ähnliche emotionale Reaktionen auslösen können wie echte Bilder. Dennoch können subtile Fehler oder unnatürliche Texturen im Bild Unbehagen oder Angst über vermeintliche Manipula-

tionsversuche auslösen, was als „uncanny valley theory“ bekannt ist (Schindler et al., 2017; Tucciarelli et al., 2022). Wenn der KI-Ursprung – auch durch fehlerhafte Bildinhalte – erkannt wird, können die Bilder als künstlich oder manipulativ bewertet werden (Appel & Prietzel, 2022; Lu & Yuan, 2024). Als Folge kann die wahrgenommene Aufdringlichkeit erhöht werden, da die Rezipient*innen sich möglicher Manipulationsversuche stärker bewusst sind, was zu Reaktanz führen kann (Vaccari & Chadwick, 2020; Weikmann et al., 2024). Im Gegensatz dazu können fortschrittliche KI-Tools sehr realitätsnahe Bilder erzeugen, die oft nicht von echten Bildern zu unterscheiden sind; was die wahrgenommene Natürlichkeit und Überzeugungskraft von KI-Botschaften steigern kann (Bai et al., 2023; Nightingale & Farid, 2022).

(2) Dabei können die emotionalen Reaktionen auf KI-generierte Bilder und ihre Bewertungen stark variieren, abhängig von der *Valenz der dargestellten Inhalte* und dem dadurch gegebenen emotionalen Rahmen (sogenanntes visuelles, emotionales Framing; Geise & Lobinger, 2015; Iyer et al., 2014; Nabi, 2003). Studienergebnisse zeigen, dass positive KI-Bilder, wie lächelnde Gesichter, Hoffnung auslösen (Eiserbeck et al., 2023) und dadurch auch als natürlicher und überzeugender empfunden werden können (Geise, 2011; Maurer, 2016). Dagegen rufen negative KI-Bilder, wie wütende oder traurige Gesichter, eher Angst oder Wut hervor (Eiserbeck et al., 2023) und können in Folge die Wahrnehmung von Künstlichkeit und Aufdringlichkeit erhöhen (Masch & Gabriel, 2020). Die Forschung von Ba-

rari et al. (2021) und Hameleers et al. (2022) zeigt zudem, dass negative Deepfakes zwar glaubwürdig sein können, aber nicht zwingend überzeugender sind als andere Formen von Fehlinformationen.

(3) Zur Rolle von *Aufklärungstexten und KI-Kennzeichnungen* für die emotionale Wirkung von KI-generierten Bildern ist bisher wenig bekannt. Studien weisen jedoch darauf hin, dass Hinweise auf den KI-Ursprung die wahrgenommene Authentizität der dargestellten Inhalte und die Glaubwürdigkeit des Bildes beeinflussen (Eiserbeck et al., 2023; Lu & Yuan, 2024). Explizite KI-Hinweise aktivieren das Persuasionswissen der Rezipient*innen, was als kognitiver Abwehrmechanismus fungiert und die Überzeugungskraft der Botschaft schwächen kann (Friestad & Wright, 1994). Lu und Yuan (2024) fanden zudem, dass solche Kennzeichnungen die Wahrnehmung der Bilder als aufdringlich verstärken können, da sie das Bewusstsein für mögliche Manipulationen schärfen und dadurch auch Emotionen wie Wut oder Angst auslösen können.

(4) Abschließend zeigen Studien, dass *Eigenschaften der Rezipient*innen* wie interne und externe politische Selbstwirksamkeit, die Einstellung zur Absenderpartei und Voreinstellungen zum Thema die Bewertung von KI-generierten Kampagnenbildern beeinflussen können. Personen mit hoher politischer Selbstwirksamkeit setzen sich intensiver mit Botschaften auseinander und sind weniger anfällig für oberflächliche Persuasionsversuche (Dobber et al., 2021; Łabuz & Nehring, 2024; Lu & Yuan, 2024). Sympathien

für die Absenderpartei führen zu einer positiveren Bewertung der Bilder, während ablehnende Haltungen oder gegensätzliche Themen Reaktanz und eine kritischere Beurteilung auslösen können (Kunda, 1990; Wu, 2024).

Auch persönliche Merkmale wie Geschlecht und Alter spielen eine Rolle: Jüngere Zielgruppen, insbesondere Männer, stehen neuen Technologien offener gegenüber und bewerten KI-generierte Inhalte oft positiver, während ältere Personen tendenziell skeptischer sind (Sundar, 2008; Wang et al., 2024). Jüngere Menschen reagieren stärker auf klimabezogene Themen, während ältere sensibler auf Migrationsthemen reagieren (Geise, 2011; Leonhard & Bartsch, 2020).

An diesen Forschungsstand wollen wir anknüpfen und für realitätsnahe KI-Bilder im Kontext von politischen Kampagnen herausfinden, welche Emotionen bei den Betrachter*innen ausgelöst werden und wie sie diese Inhalte bewerten.

Fragestellung 8: Welche Wirkung geht von der Rezeption KI-generierter Kampagnenbilder aus?

Neben den Fragen zur Wahrnehmung KI-generierter Kampagnenbilder beschäftigen wir uns auch damit, wie sich die Rezeption von KI-Botschaften auf die Einstellungen der Betrachter*innen auswirken. Dabei interessiert uns insbesondere die allgemeine Akzeptanz des Einsatzes von KI in der politischen Kommunikation, die Bewertung von Parteien, die KI in ihrer Arbeit einsetzen und die Bewertung von KI als Gefahr für die Demokratie. Einen ersten Anhaltspunkt bieten die Ergebnisse unserer repräsentativen Befragung. Diese

hat gezeigt, dass die Einstellungen zum Einsatz von KI in der Politik gespalten sind. Steht einerseits ein großer Anteil dem Einsatz skeptisch gegenüber, empfinden andere die neuen technologischen Möglichkeiten als nützlich. Inwiefern diese unterschiedlichen Bewertungen mit den konkreten Eigenschaften der generierten Inhalte zusammenhängen, ist bislang noch nicht wissenschaftlich erforscht worden. Um diese Forschungslücke auszuleuchten, untersuchen wir in unserem Experiment die Rolle von Bildinhalten, der Bildwahrnehmung und der Eigenschaften der Rezipient*innen. Dabei stehen die folgenden Forschungsfragen zu den zu beeinflussenden Einstellungen im Mittelpunkt des Forschungsinteresses:

Wie beeinflusst die Rezeption KI-generierter Kampagnenbilder die Bewertung des Einsatzes von KI in Kampagnen (Fragestellung 8.1), die Bewertung von Parteien, die KI einsetzen (Fragestellung 8.2) und die Bewertung von KI als Gefahr (Fragestellung 8.3)?

Die Wirkung von KI-generierten Bildern in politischen Kampagnen kann durch verschiedene Faktoren beeinflusst werden, darunter die emotionale Valenz der Darstellungen, die thematische Übereinstimmung mit den Einstellungen der Rezipient*innen und die Authentizität der Inhalte. Entlang dieser Variablen skizzieren wir im Folgenden den aktuellen Forschungsstand zum Einfluss von KI-generierten Bildinhalten auf die Einstellungen zum Einsatz von KI in Kampagnen:

KI-Bildinhalte (emotionale Valenz; Einstellungskongruenz): Die Wahrnehmung und Bewertung

einzelner KI-generierter Inhalte kann die generelle Wahrnehmung und Akzeptanz von KI-Technologien in politischen Kampagnen beeinflussen. Zu den wichtigsten Einflussfaktoren auf die mögliche Wirkung von Kampagneninhalten zählen insbesondere die Valenz und das dargestellte Thema.

Die Forschung zur affektiven Informationsverarbeitung zeigt, dass die Rezeption und Verarbeitung von positiven oder negativen Medieninhalten zu entsprechenden kurzfristigen Emotionen führen kann und letztlich Einstellungen zum Medium und zu den Absender*innen beeinflussen können (z.B. Kühne et al., 2012; Maurer 2014; Schemer 2009). So konnte beispielsweise gezeigt werden, dass emotionalisierende Gestaltungsmittel in politischen Kampagnenbotschaften – darunter Angriffe auf Kandidierende, stimmungsgeladene Bildsprache oder Themen – das emotionale Erleben der Rezipient*innen prägen und damit sowohl negative als auch positive Einstellungen hervorrufen können (Brader 2006; Scheufele & Gasteiger 2007; Thorson et al., 1991).

Nach dem „Affect Infusion Model“ (Forgas, 1995) gibt es in Anlehnung an Zwei-Prozess-Modelle der Informationsverarbeitung zwei zentrale Mechanismen: Bei heuristischer, oberflächlicher Verarbeitung dient die Stimmung direkt als Bewertungsheuristik („affect-as-information“). Bei intensiver, substanzieller Verarbeitung wirkt die Stimmung dagegen indirekt über „affektives Priming“ – sie beeinflusst, welche Informationen bei der Bewertung stärker berücksichtigt und wie

sie interpretiert werden. So signalisieren positive Gefühle, dass die Situation unproblematisch ist, was zu einer eher heuristischen, intuitiven Verarbeitungsweise führt. Daher können positive KI-Kampagneninhalte zu einer weniger kritischen Bewertung der Technologie und ihrer Anwender*innen führen – denn die Rezipient*innen sind in positiver Stimmung eher bereit, neue Entwicklungen zu akzeptieren und Risiken einzugehen. Negative Gefühle hingegen aktivieren eine systematischere und analytischere Verarbeitung, da sie dem Individuum signalisieren, dass die Situation möglicherweise problematisch ist (Fiedler, 1988). Negative KI-Kampagneninhalte könnten daher eine intensivere kritische Auseinandersetzung mit möglichen Risiken der Technologie anstoßen.

Darüber hinaus spielt die thematische Übereinstimmung der Bilder mit den Einstellungen der Rezipient*innen eine wesentliche Rolle (vgl. Kunda, 1990; Wu, 2024). So sollten Bilder, die Themen auf eine Art darstellen, die im Einklang mit den Überzeugungen der Betrachter*innen stehen, zu einer positiven Einstellung gegenüber KI führen, während Bilder, die stark von den eigenen Überzeugungen abweichen, die Einschätzung von KI als Gefahr verstärken können. Das liegt daran, dass Inhalte, die nicht mit den eigenen Voreinstellungen übereinstimmen, das kognitive Dissonanzempfinden erhöhen (Festinger, 1957; Kruglanski, 1990). Da Menschen bestrebt sind, Konsonanz herzustellen, versuchen sie zukünftig, Dissonanzquellen zu vermeiden oder die Quelle solcher Inhalte abzuwerten. Dieses Erkenntnis wird durch eine Studie von Dobber

et al. (2021) gestützt, die feststellten, dass die Wirkung von KI-generierten Deepfakes von der Übereinstimmung mit den eigenen politischen Ansichten abhängt. Wähler*innen, die mit den dargestellten Ansichten übereinstimmten, zeigten eine stärkere positive Reaktion auf die Deepfake-Botschaften, was sich schließlich in einer grundlegend positiveren Einstellung zu KI-generierten Inhalten niederschlagen könnte.

Bildwahrnehmung (emotionale Reaktion; Bildbewertung): KI-generierte Bildinhalte können durch ihre Gestaltung und Präsentation unterschiedliche emotionale Reaktionen hervorrufen, die die Verarbeitung von Informationen und letztlich auch die Meinungsbildung beeinflussen (Geise, 2011; Scheufele & Gasteiger, 2007; Schemer, 2014). Aus psychologischer Perspektive spielen affektive Reaktionen eine zentrale Rolle in der Informationsverarbeitung, da Emotionen als Heuristiken fungieren, die die kognitive Belastung reduzieren und schnelle Bewertungen ermöglichen (Forgas, 1995). Insbesondere Emotionen wie Wut, Angst oder Hoffnung aktivieren spezifische Pfade der Informationsverarbeitung, die entweder verstärkte Aufmerksamkeit oder Vermeidungsstrategien auslösen können (Marcus et al., 2000). Studien zeigen, dass solche Emotionen Auswirkungen auf die Einstellung zu Absender*innen oder thematisierten Inhalten haben können und die Bereitschaft zur Verhaltensänderung oder zur Unterstützung ihnen gegenüber erhöhen (Geise, 2011; Leonhard & Bartsch, 2022; Schemer, 2014). Dabei ist besonders hervorzuheben, dass visuelle Informationen durch emotionale Reaktionen oft stärker

wirken als durch Sachinhalte vermittelte (Kepplinger & Maurer, 2005). Daher kann die bei der Rezeption von KI-generierten Inhalten empfundene Wut, Angst oder Hoffnung einen negativen beziehungsweise positiven Einfluss auf die Bewertung des Einsatzes von KI in Kampagnen und auf Parteien haben, die KI einsetzen (vgl. Vaccari & Chadwick, 2020).

Darüber hinaus ist die Bewertung von KI-generierten Bildern als unproblematisch, authentisch oder natürlich ein weiterer Faktor für den möglichen Einfluss auf Einstellungen zu KI oder ihrer Nutzer*innen. Im Lichte von Zwei-Prozess-Modellen wie dem „Elaboration Likelihood Model“ (ELM) von Petty und Cacioppo (1986) erfolgt die Informationsverarbeitung bei der Wahrnehmung eines Bildes als authentisch oder natürlich eher über die zentrale Route, bei der sich Rezipient*innen intensiv mit den Inhalten auseinandersetzen. Authentizität und Natürlichkeit fördern Glaubwürdigkeit und vertrauenswürdige Absichten seitens der Urheber*innen, was die Bereitschaft erhöht, die Botschaft aufzunehmen und sich mit ihr zu beschäftigen (Geise, 2011; Scheufele & Gasteiger, 2007; Schemer, 2014). Beim Einsatz von KI in politischen Kampagnen kann dies die positive Wahrnehmung des Einsatzes von KI und ihrer Absender*innen fördern (vgl. Dobber et al. 2021). Im Gegensatz dazu verstärken Bilder mit einem erkennbar künstlichen oder unechten Erscheinungsbild Skepsis, da sie Zweifel an der Intention und Integrität des Urhebers schüren (Hameleers et al., 2024). Besonders kritisch sind Deepfakes, die aufgrund ihres Potenzials zur Täuschung oft Unsicherheit

auslösen und das Vertrauen in die Informationsquelle reduzieren können. Dies wirkt sich nicht nur negativ auf die Glaubwürdigkeit der Inhalte aus, sondern auch auf die allgemeine Akzeptanz der dahinterstehenden Technologie.

*Eigenschaften von Rezipient*innen (interne und externe politische Selbstwirksamkeit; KI-Wissen/Kompetenz; Alter; Geschlecht):* Die individuellen Eigenschaften der Rezipient*innen, wie politische Selbstwirksamkeit, KI-Wissen und demografische Merkmale, haben ebenfalls einen bedeutenden Einfluss darauf, wie der Einsatz von KI in politischen Kampagnen wahrgenommen und bewertet wird. Politische Selbstwirksamkeit, sowohl intern (eigene Kompetenzüberzeugung) als auch extern (Vertrauen in Problemlösungsfähigkeit des politischen Systems), beeinflusst die Bereitschaft der Betrachter*innen, sich intensiv mit den Kampagnenbotschaften auseinanderzusetzen und sie als relevant für ihre politische Partizipation zu betrachten (Lu & Yuan, 2024). Personen mit hoher politischer Selbstwirksamkeit fühlen sich eher in der Lage, die durch KI-generierte Inhalte vermittelten Botschaften kritisch zu hinterfragen und ihre Relevanz für ihre eigene politische Entscheidungsfindung zu erkennen. Diese Erkenntnis wird durch die Studie von Dobber et al. (2021) gestützt, die zeigte, dass Wähler*innen, die sich stärker mit einer Partei identifizierten, weniger anfällig für die negativen Effekte von KI-generierten Deepfakes waren.

Ein umfassendes Wissen über KI kann dazu führen, dass Rezipient*innen Inhalte analytischer und kritischer bewerten, was zu einer differen-

zierteren und realistischeren Einschätzung der Risiken und Chancen von KI führt (Appel & Prietzel, 2022; Wang et al., 2024). Eiserbeck et al. (2023) bestätigten diese Annahme in ihrer Studie, in der sie feststellten, dass Teilnehmer*innen mit höheren Werten in einem kognitiven Reflexionstest KI-generierte Inhalte als weniger glaubwürdig einschätzten. Im Gegensatz dazu könnten weniger informierte Rezipient*innen stärker auf emotionale Reize reagieren, was zu einer weniger reflektierten und potenziell vor-eingenommenen Einstellung gegenüber KI in politischen Kampagnen führen könnte. Darüber hinaus zeigen Studien, dass demografische Faktoren wie Alter und Geschlecht ebenfalls die Einstellungen zu KI beeinflussen. Jüngere Zielgruppen, insbesondere Männer, zeigen oft eine größere Offenheit für neue Technologien und haben tendenziell eine geringere Risikoeinschätzung gegenüber KI in politischen Kampagnen, während ältere Zielgruppen möglicherweise konservativere und skeptischere Haltungen einnehmen (Sundar, 2008; Wang et al., 2024).

4.2 Methode und Design des Experiments

Um die aufgeworfenen Fragen zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern zu beantworten, haben wir ein Online-Experiment durchgeführt. Neben dem Hinzufügen von Vorab-Hinweisen auf KI und Disclaimern im Bild wurden außerdem die Valenz der Bilder sowie die Thematik variiert.

Diese Variationen haben wir mit Hilfe generativer KI vorgenommen: wir haben eine Kombination

aus von Midjourney erstellten Bildern und mit ChatGPT erstellten Slogans als Stimuli genutzt. Wie wir dabei vorgegangen sind, erläutern wir im Folgenden.

4.2.1 Experimentelle Variation(en) durch Prompting

Für die in unserer Studie verwendeten Stimuli (siehe Infobox 4) nutzten wir die Strategie des Political Prompt Engineering (siehe Kapitel 2.1.2; Infobox 2). Hierfür verwendeten wir das TTI-Modell Midjourney. Unabhängig von dem Thema des zu erstellenden Bildes (Migration, Umwelt und Agrar), seiner politischen Richtung (progressiv und konservativ), dem Absender (reale Partei: Grüne oder CDU, fiktive Initiative: Initiative360 oder BürgerVision) und der über das Bild vermittelten Stimmung (positiv und negativ) waren die Prompts in ihrem Aufbau allesamt gleich gestaltet und lehnten sich an das Prompt Engineering für realistische Bilderstellungen an (siehe als Beispiel Abbildung 1).⁴ So wurden die für das Bild relevanten Akteur*innen in der Regel zentral im Vordergrund dargestellt und bildeten dabei zugleich die Stimmung ab, die durch das Bild transportiert werden sollte. Weitere Umgebungsfaktoren waren rund um den jeweiligen Bildakteur im Hintergrund gruppiert und entsprechend dem Bildthema gewählt. Durch zusätzliche Adjektive wurde die vermittelte Stimmung des zu erstellenden Bildes spezifiziert. Weitere stimmungstechnische Elemente

wurden durch die Kamerawinkel und die im Bild zu sehende Beleuchtung sowie Midjourney-spezifische Parameter bestimmt. Für jedes Motiv (bspw. überschwemmtes Dorf) wurden jeweils zwei KI-Bilder erstellt. Um im weiteren Verlauf den Effekt der KI-Bilder zu testen, wurde zusätzlich zu der zweifachen Ausführung der KI-Bilder eines Motivs, ein reales Stockfoto hinzugefügt, das den KI-generierten in ihrer Gestaltung möglichst genau entspricht.

Die Bilder wurden anschließend im Bildbearbeitungsprogramm Photoshop so angepasst, dass sie der Ästhetik typischer Wahlkampfplakate ähneln. Hierfür wurden die Bilder je nach Thema, politischer Richtung und vermittelter Stimmung mit weiteren visuellen Elementen versehen. Dies umfasste zum einen die Logos von jeweils einem real existierenden politischen Absender aus dem konservativen (CDU) und dem progressiven Spektrum (Bündnis 90/Die Grünen) sowie die Logos von jeweils einem fiktiven Absender aus diesen politischen Spektren (BürgerVision; Initiative 360). Je nach Stimmung und Thema wurde jedes Bild zudem mit einer politischen Botschaft in Form einer Headline sowie einem dazugehörigen Untertitel ausgestattet. Mithilfe von ChatGPT wurden für jedes Thema zwei verschiedene Textelemente pro politischer Ausrichtung erstellt. Diese variierten nicht über die Absender der jeweiligen politischen Richtung. Letztlich wurde auch ein Disclaimer, dass es sich

4 Aufgrund der möglichen Risiken im politischen und gesellschaftlichen Kontext, die eine detaillierte Beschreibung der von uns verwendeten Prompts nach sich ziehen könnte, wird ihr Aufbau nur allgemein beschrieben. Auf eine Ausführung konkreter Beispiele zur Erstellung der Stimuli-Materialien für die Studie wird verzichtet.



Infobox 4:

Beispiele der KI-generierten Stimuli

Stimulus 1: Umwelt, negativ, progressiv



Stimulus 2: Umwelt, positiv, konservativ



Stimulus 3: Migration, negativ, konservativ



Stimulus 4: Migration, positiv, progressiv



Stimulus 5: Agrar, negativ, konservativ



Stimulus 6: Agrar, positiv, progressiv



bei dem entsprechenden Bild um ein KI-generiertes handelte, hinzugefügt. Dieser lehnte sich an die Kennzeichnung von KI-Inhalten von Facebook und Instagram an (siehe Abbildung 9; Meta, 2024).

4.2.2 Stichprobe und Randomisierung

Für die experimentelle Befragung zur Wahrnehmung und Wirkung von KI-generierten Inhalten wurden parallel zur repräsentativen Befragung weitere Teilnehmende durch

den Anbieter Cint rekrutiert. Wie in der Befragungsstudie wurde eine hinsichtlich Alter und Geschlecht repräsentative Stichprobe der deutschen Bevölkerung zwischen 18 und 69 Jahre angestrebt. Insgesamt nahmen 3.340 Personen an der Befragung im Zeitraum vom 7. bis 25. Mai 2024 teil.

Diese Stichprobe wurde anschließend durch Qualitätschecks (Straightlining, Wissensfragen) bereinigt. Nach der Bereinigung wurden 2.081 Teilnehmende berücksichtigt. Auch diese Stichprobe weicht leicht von der angestrebten Quotierung ab. Da uns im Experiment zur Wirkung und Wahrnehmung KI generierter Inhalte jedoch keine absoluten Werte, sondern Zusammenhänge interessieren, ist eine Gewichtung nicht notwendig.

Für die Stichprobe ergibt sich ein Anteil von rund 53 Prozent weiblicher Befragter sowie 47 Prozent männlicher Befragter; vier Befragte identifizierten sich als „divers“. Wie bei der repräsentativen Umfrage lag das Durchschnittsalter der Befragten bei 44 Jahren, wobei knapp 30 Prozent der Altersgruppe zwischen 30 und 39 Jahren angehörten. Auf die Gruppen 40 und 49 sowie 50 und 59 Jahren entfielen jeweils rund ein Viertel. Die Altersgruppe 18 bis 29 Jahre machte rund 10 Prozent, die Altersgruppe zwischen 60 und 69 Jahren machte rund 12 Prozent aus. Die Anteile der Beschäftigungsverhältnisse und Bildungsabschlüsse weichen nur geringfügig von denen der Befragten in der repräsentativen Umfrage ab.

4.2.3 Untersuchungsdesign und -ablauf

Die von uns verfolgten Forschungsziele haben wir mithilfe eines aufwendigen methodischen Designs umgesetzt (siehe Abbildung 20). Zunächst haben wir die politischen Einstellungen mit jeweils drei Items gemessen (5er-Skalen, 1 = „stimme überhaupt nicht zu“ bis 5 = „stimme voll und ganz zu“). Erhoben haben wir das Interesse an Politik (5er-Skala, 1 = „überhaupt nicht interessiert“ bis 5 = „sehr stark“) sowie die Einschätzung politischer Selbstwirksamkeit, wobei wir zwischen der internen (z. B. „Wichtige politische Fragen kann ich gut verstehen und einschätzen“) und externen (z. B. „Die Politiker kümmern sich darum, was einfache Leute denken“) Einschätzung unterschieden haben. Zudem sollten die Teilnehmenden angeben, wie häufig sie sich über die klassischen Massenmedien sowie über die digitalen Plattformen (z. B. Instagram, YouTube) informieren (7er-Skala; 0 = „nie“ bis 6 = „mehrmals täglich“). Zudem haben wir die politische Nutzung (z. B. „Ich kommentiere Beiträge von Politikern oder Parteien“) der am intensivsten genutzten Plattform erfragt (6 Items, 5er-Skalen, 1 = „nie“ bis 6 = „sehr häufig“). Mit jeweils drei Items wurden die Voreinstellungen zu den Themen Klimaschutz (z. B. „Klimawandel ist ein großes Problem“), Migration (z. B. „Immigration sollte weiter begrenzt werden“) und Landwirtschaft (z. B. „Durch die Massenproduktion in der Landwirtschaft geht die Wertschätzung für das einzelne Tier verloren“) abgefragt (5er-Skalen, 1 = „stimme überhaupt nicht zu“ bis 5 = „stimme voll und ganz zu“). Diese Variablen können einen Einfluss auf

Abbildung 20:

Untersuchungsdesign und -ablauf des Online-Experiments zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern

1. Frageblock

- **Soziodemographie**
- **Politische Einstellungen** (Wichtigstes Thema, politisches Interesse, politische Wirksamkeit, Parteineigung, Voreinstellungen)
- **Mediennutzung** (Social-Media-Nutzung, Politische Social-Media-Nutzung)

Experimentalteil

randomisiert

Zufällige Einblendung: Aufklärungstext zur KI-Nutzung

Hinweis, dass folgende
Werbeanzeigen mit KI
generiert sein können

randomisiert

Einblendung von 3 Stimuli

Zufällige Variation von
6 Bildeigenschaften:

- *Bildmotiv* (KI-generiert vs. echtes Bild)
- *Thema* (Migration, Umwelt, Agrar)
- *Politische Richtung* (progressiv vs. konservativ)
- *Absender* (reale Partei vs. fiktive Initiative)
- *Stimmung* (positiv vs. negativ)
- *KI-Kennzeichnung* (angezeigt vs. nicht angezeigt)



3. Frageblock: Wirkung

- Bewertung von KI in der Politik
- Konsequenzen von KI in der Politik
- Einsatzgebiete von KI in der Politik
- Einschätzung von Parteien, die KI einsetzen

2. Frageblock: Wahrnehmung (des letzten Bildes)

- Erkennung KI-Generierung
- Evozierte Emotionen
- Bildwahrnehmung

Quelle: Eigene Darstellung.

die Erkennung und Bewertung von KI-generierten Inhalten haben.

Anschließend wurden die Befragten in zwei Gruppen eingeteilt. Eine Gruppe sah einen Hinweistext, bei dem lediglich angekündigt wurde, dass sie im Folgenden drei politische Anzeigen gezeigt bekämen; die andere Hälfte bekam die gleiche Ankündigung, ergänzt um einen Hinweis auf potenzielle KI-Bilder („Manche der verwendeten

Bilder könnten mit Hilfe künstlicher Intelligenz erstellt worden sein und bilden keine realen Personen, Situationen oder Gegenstände ab“). Im Anschluss wurden allen Befragten nacheinander drei Bilder gezeigt. Dabei wurden die Befragten erneut zufällig einer Gruppe zugewiesen: der einen Hälfte wurden alle drei Bilder versehen mit einem Hinweis, dass es sich um ein KI-generiertes Bild handelt, präsentiert; die andere Hälfte sah alle Bilder ohne einen solchen Hinweis.

Im Anschluss wurden die Teilnehmenden für das letzte gezeigte Bild gefragt, in welchem Ausmaß die Emotionen Wut, Angst und Hoffnung durch das Bild ausgelöst wurden (7er-Skala, 1 = „überhaupt nicht“ bis 7 = „in sehr großem Ausmaß“). Überdies wurde die Wahrnehmung der Anzeigen auf semantischen Differenzialen (z. B. „künstlich“ vs. „natürlich“) erfragt (9 Items, 7er-Skalen). Die möglichen Einflüsse der von uns untersuchten Faktoren (KI-generiert, Hinweistext, Texthinweis im Bild, Valenz) werden in Mittelwertunterschieden dargestellt.⁵ Bei der Schätzung der Mittelwerte haben wir die zuvor erhobenen Einflussfaktoren kontrolliert.

4.3 Ergebnisse des Experiments

Im Folgenden stellen wir die Ergebnisse unseres Experiments zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern vor. In einem ersten Schritt diskutieren wir die Fähigkeit zur Erkennung der KI-Kampagnenbilder, danach die durch das Betrachten der Bilder ausgelösten emotionalen Reaktionen und die Bewertung der Botschaften hinsichtlich Überzeugungskraft, Natürlichkeit und Aufdringlichkeit. In einem zweiten Schritt geht es um den Einfluss, den die Konfrontation mit KI-generierten Kampagneninhalten auf generelle Einstellungen der Bürger*innen zum Einsatz von KI in Kampagnen und die Bewertung von Parteien, die KI einsetzen, entfalten kann. Dabei berücksichtigen wir individuelle Eigenschaften der Befragten und greifen ebenso

auf die Bewertungen der Befragten zurück, die sie nach der Betrachtung aller drei gezeigten Kampagnenbotschaften in unserem Experiment gegeben haben.

4.3.1 Erkennung von KI-generierten Kampagnenbildern

In unserer Untersuchung zur Wahrnehmung von KI-generierten Bildern haben wir zunächst untersucht, welche Faktoren die Erkennung von KI-Bildern beeinflussen. Dazu haben wir die Teilnehmenden gefragt, ob sie glauben, dass das Bild, welches sie zuletzt gesehen haben, mit KI erstellt wurde. Dabei konnten die Befragten ihre Einschätzung auf einer Skala einordnen (1 = „auf keinen Fall“ bis 5 = „auf jeden Fall“). Zunächst lässt sich generell festhalten, dass die Befragten über alle Variationen der Bildinhalte hinweg Schwierigkeiten bei der Erkennung des KI-generierten Ursprungs haben und dabei jeweils im Mittelfeld unserer Wahrscheinlichkeitsskala landen (siehe Abbildung 21). Unabhängig davon, ob das Bild tatsächlich KI-generiert war oder nicht, wurden sowohl reale als auch KI-generierte Bilder mit nahezu gleicher Wahrscheinlichkeit als KI-generiert eingeschätzt (siehe Abbildung 21 links). Obwohl mit KI erstellte Bilder etwas häufiger als solche erkannt werden, ist dieser Unterschied minimal, bedenkt man die Spannbreite der Skala. Auch zeigt sich, dass positivere Bilder weniger wahrscheinlich als KI-generiert wahrgenommen werden (siehe Abbildung 21 Mitte links). Dies deutet darauf hin, dass Menschen

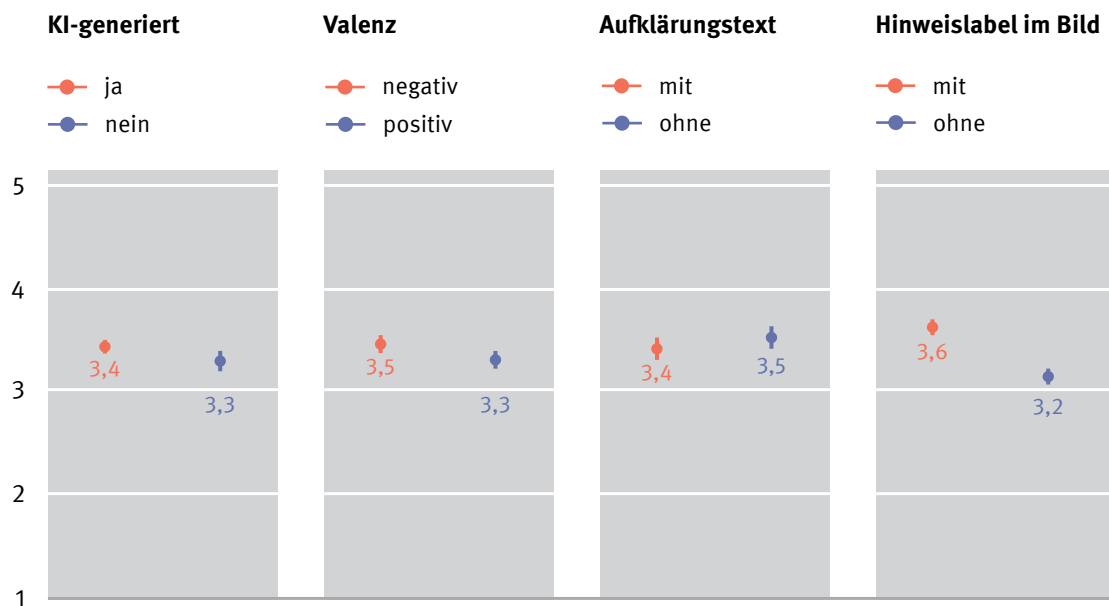
⁵ Hier soll noch einmal darauf hingewiesen werden, dass wir auch die Botschaftstexte entsprechend der im Bild gezeigten Emotionen (positiv/negativ) und Themen (Migration/Umwelt/Agrarpolitik) angepasst haben, sodass diese ebenfalls einen Einfluss auf die Wahrnehmung haben können.

dazu neigen, KI-generierte Inhalte eher mit neutralen oder negativen Darstellungen zu verknüpfen, während sie positivere Bilder eher als authentisch einschätzen, wenngleich der Unterschied auch hier minimal ist. Die Kennzeichnung der Bilder mit einem entsprechenden grafischen Hinweis im Bild selbst hat hingegen einen substantziellen Einfluss auf die Wahrnehmung von KI (siehe Abbildung 21 rechts). Hier unterscheidet sich die Wahrnehmung immerhin um einen halben Skalenpunkt. Bemerkenswert ist der Befund zum Einfluss unseres Aufklärungstextes, der Versuchsteilnehmer*innen in einer

der Experimentalgruppen darauf hinwies, dass manche genutzten Bilder im Experiment mit KI erstellt wurden und keine realen Personen, Situationen oder Gegenstände abbilden: Dieser Hinweis senkte die Wahrscheinlichkeit, die gezeigten Bilder als KI-generiert einzustufen, minimal (siehe Abbildung 21 Mitte rechts). Das unterstreicht, wie wichtig klare Kennzeichnungen sind, um die Transparenz von KI-Bildern zu erhöhen und dass eine allgemeine Aufklärung über den Einsatz von KI in der politischen Kommunikation nicht automatisch zu einer erhöhten Erkennung führt.

Abbildung 21:

Erkennung KI-generierter Bilder abhängig von tatsächlicher KI-Generierung, Valenz, Aufklärungstext und Hinweislabel



Fragestext: „Bitte erinnern Sie sich nochmal an das Bild zurück, das Sie als letztes gesehen haben. Was glauben Sie, wurde das Bild mit Hilfe künstlicher Intelligenz (KI) erstellt?“ (1 = „auf keinen Fall“ bis 5 = „auf jeden Fall“).

Quelle: Eigene Darstellung.

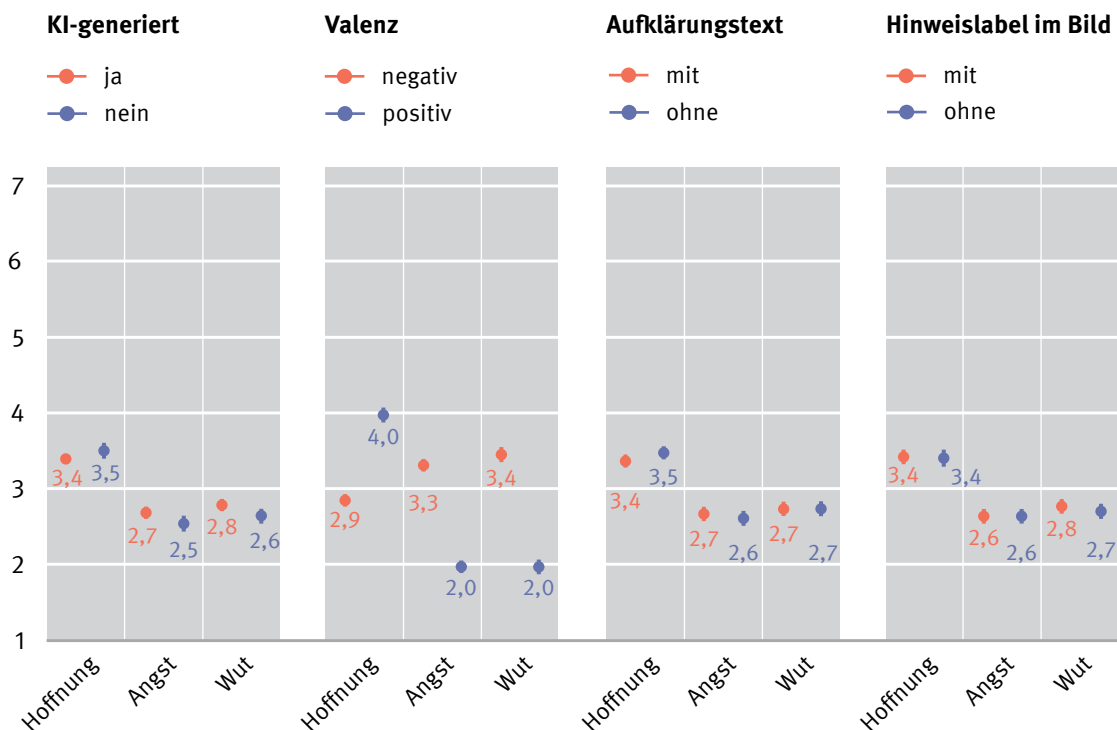
4.3.2 Ausgelöste Emotionen durch KI-generierte Kampagnenbilder

In einem zweiten Schritt haben wir Faktoren analysiert, die beeinflussen können, ob KI-generierte Bilder bei den Betrachter*innen Hoffnung, Angst oder Wut auslösen. Nach jedem gezeigten Bild haben wir die Befragten gebeten, anzugeben, in welchem Ausmaß sie die jeweiligen Emotionen verspüren. Ihre Einschätzung konnten die Teilnehmenden auf einer Skala angeben (1 = „überhaupt nicht“ bis 7 = „in sehr großem Maße“).

Zunächst lässt sich für die Bildinhalte feststellen, dass sich KI-generierte Bilder in ihrer Wirkung kaum von echten Fotografien unterscheiden und nur minimal weniger Hoffnung, mehr Angst oder Wut hervorrufen als nicht KI-generierte Bilder (siehe Abbildung 22 links). Den entscheidenden Unterschied macht die Valenz: Positivere Bilder lösen deutlich mehr Hoffnung aus und signifikant weniger Angst sowie Wut aus (siehe Abbildung 22 Mitte links). Dies unterstreicht die Bedeutung des Bildinhalts für die emotionale Reaktion – unabhängig davon,

Abbildung 22:

Durch die Kampagnenbilder ausgelöste Emotionen abhängig von tatsächlicher KI-Generierung, Valenz, Aufklärungstext und Hinweislabel im Bild



Fragetext: „Wenn Sie an das Bild denken, das Sie eben gesehen haben: In welchem Ausmaß löst das Bild folgende Gefühle und Empfindungen bei Ihnen aus?“ (1 = „überhaupt nicht“ bis 7 = „in sehr großem Maße“).

Quelle: Eigene Darstellung.

ob das Bild künstlich oder echt ist. Weder KI-Disclaimer im Bild noch der Aufklärungstext über KI-generierte Bilder in politischen Kampagnen haben Einfluss auf die emotionalen Reaktionen (siehe Abbildung 22 Mitte rechts und rechts).

Interessanterweise zeigen sich deutliche Unterschiede in den emotionalen Reaktionen je nach angesprochenen Themen in den Bildern. Im Vergleich zu Agrar-Themen lösen Migrationsthemen weniger Hoffnung, aber mehr Angst und Wut aus. Umweltthemen führen im Vergleich zu Agrar-The-

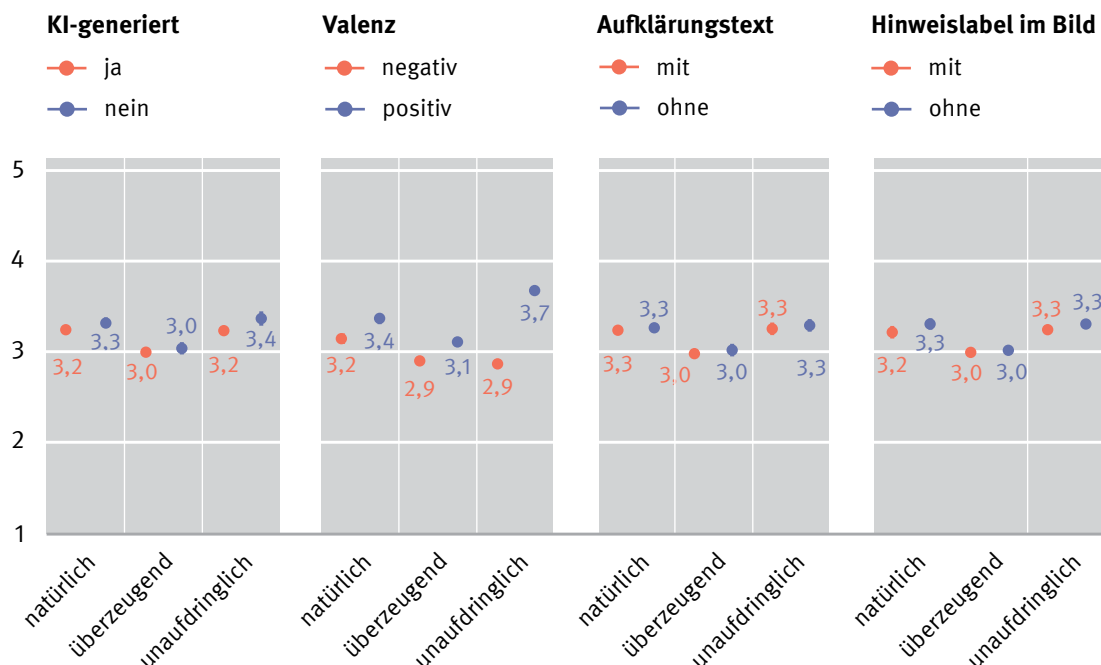
men ebenfalls zu weniger Hoffnung sowie zu mehr Angst ($b = 0.28, p < 0.001$) und mehr Wut, wobei der Effekt weniger stark ausgeprägt ist als bei Migrationsthemen. Die emotionale Reaktion scheint somit eher vom Bildinhalt als von der Kennzeichnung abzuhängen.

4.3.3 Bewertung der KI-generierten Kampagnenbilder

In einem dritten Schritt haben wir uns angeschaut, welche Faktoren beeinflussen, ob KI-generierte Bilder als unaufdringlich, überzeugend beziehungsweise persuasiv oder natürlich be-

Abbildung 23:

Bewertung der Kampagnenbilder in ihrer Natürlichkeit, Überzeugungskraft und Aufdringlichkeit abhängig von tatsächlicher KI-Generierung, Valenz, Aufklärungstext und Hinweislabel



Fragetext: „Wenn Sie nochmal an das gerade gezeigte Bild zurückdenken, wie haben Sie dieses wahrgenommen? Das Bild war...“ (Semantisches Differential 1-5; „natürlich“ vs „unnatürlich“, „überzeugend“ vs „nicht überzeugend“, „unaufdringlich“ vs „aufdringlich“).

Quelle: Eigene Darstellung.

wertet werden. Alle Einschätzungen wurden für jedes der drei gezeigten Bilder auf 5-stufigen Skalen auf semantischen Differenzialen erhoben („natürlich“ vs. „unnatürlich“, „überzeugend“ vs. „nicht überzeugend“, „unaufdringlich“ vs. „aufdringlich“). Dabei lässt sich für die Bildinhalte feststellen, dass KI-generierte Bilder im Vergleich zu nicht KI-generierten Bildern nur geringfügig aufdringlicher sowie etwas weniger überzeugend und natürlicher wahrgenommen werden (siehe Abbildung 23 links). Einen nennenswerten Einfluss auf die Beurteilung hat hingegen die Valenz der Bilder: Positivere Darstellungen werden durchweg als unaufdringlicher, überzeugender und authentischer wahrgenommen – unabhängig davon, ob sie mit KI erstellt wurden oder nicht (siehe Abbildung 23 Mitte links). Bemerkenswert ist, dass der KI-Disclaimer die Wahrnehmung der Unaufdringlichkeit und Überzeugungskraft nicht signifikant beeinflusst (siehe Abbildung 23 rechts). Ein vorangestellter Aufklärungstext über den möglichen KI-Einsatz in politischen Kampagnen lässt die Bilder minimal aufdringlicher erscheinen, während er keinen signifikanten Einfluss auf die Überzeugungskraft oder Authentizität hat (siehe Abbildung 23 Mitte rechts). Wie schon bei den ausgelösten Emotionen haben die Themen einen stärkeren Einfluss: Im Vergleich zu Agrar-Themen werden Migrationsthemen als aufdringlicher und weniger überzeugend eingestuft, während Umweltthemen in allen drei Dimensionen schwächer bewertet werden.

4.3.4 Einfluss der Valenz von KI-Bildern auf die Einstellungen zum Einsatz von KI

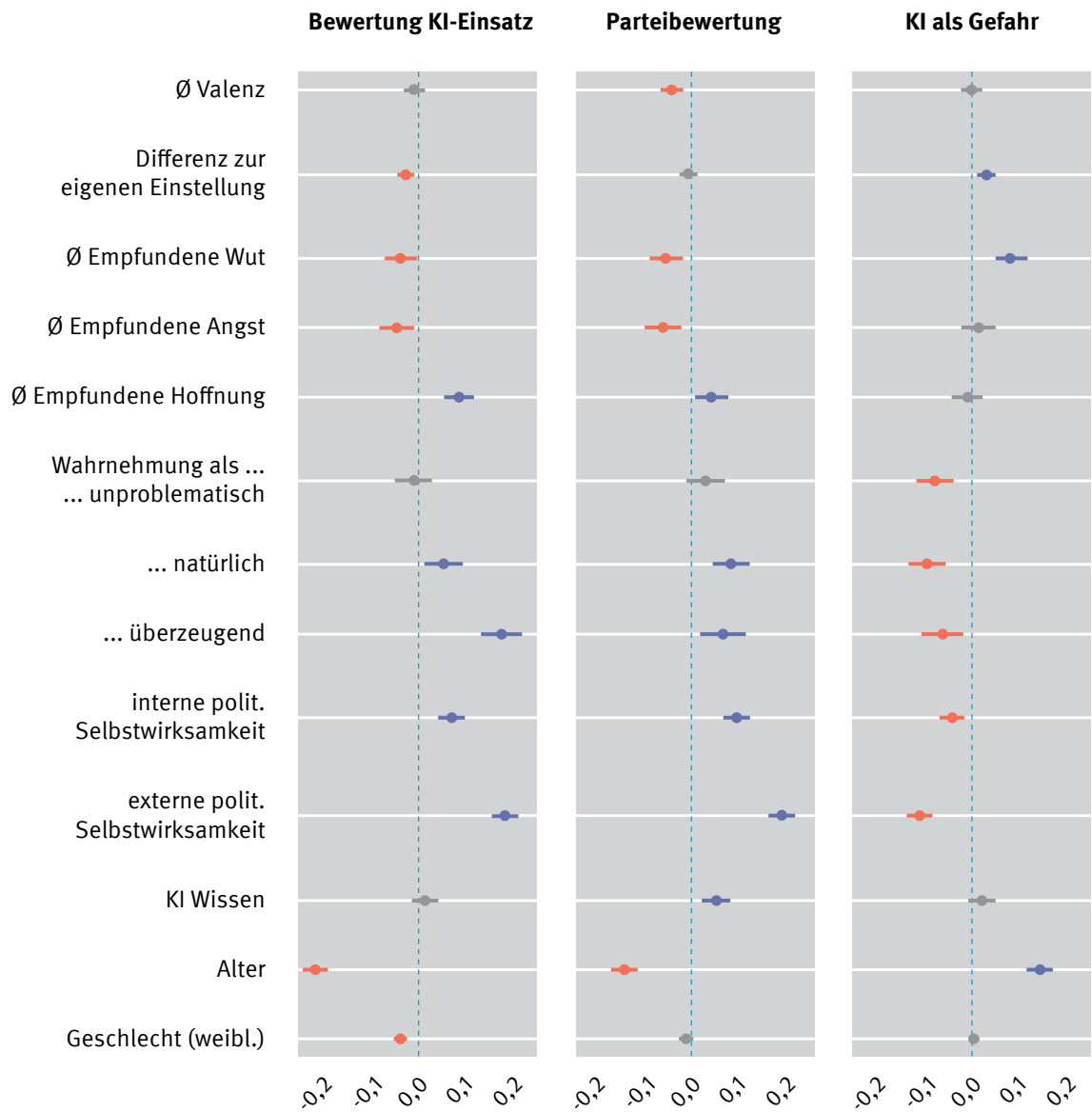
Die Ergebnisse zeigen, dass die durchschnittliche Valenz der gesehenen KI-generierten Bilder

keinen signifikanten Einfluss darauf hat, wie die Teilnehmenden den KI-Einsatz insgesamt bewerten oder ob sie KI als Gefahr wahrnehmen (siehe Abbildung 24). Bemerkenswert ist jedoch, dass Befragte, die vermehrt positive Inhalte sahen, kritischer gegenüber Parteien sind, die KI einsetzen. Dies deutet auf eine mögliche Skepsis gegenüber Parteien hin, die KI nutzen, um besonders schöne oder optimistische Bilder zu erzeugen. Darüber hinaus können wir feststellen, dass Befragte, die vermehrt Inhalte gesehen haben, die von der eigenen politischen Einstellung abwichen, den Einsatz von KI negativer bewerteten und auch eher mit Gefahr verbinden.

4.3.5 Einfluss der Wahrnehmung von KI-Bildern auf die Einstellungen zum Einsatz von KI

Es zeigt sich, dass die durchschnittlich empfundene Hoffnung nach der Betrachtung der Anzeigen positiv mit der Bewertung des KI-Einsatzes und der Einschätzung von Parteien, die KI verwenden, zusammenhängt (siehe Abbildung 24). Hingegen führt die durchschnittlich empfundene Wut dazu, dass Befragte KI als Gefahr wahrnehmen und auch Parteien, die KI einsetzen, negativer bewerten. Zudem führt die bei der Rezeption der Anzeigen ausgelöste Angst zu einer negativeren Bewertung der Parteien. Wenn die KI-Bilder überwiegend als natürlich wahrgenommen werden, bewerten die Befragten Parteien, die KI einsetzen, tendenziell positiver und sehen KI insgesamt weniger als Gefahr. Ähnlich führt die Wahrnehmung der KI-Bilder als unproblematisch dazu, dass KI weniger als Gefahr eingeschätzt wird. Darüber hinaus lassen sich positive Zusammenhänge für die durchschnittlich wahr-

Abbildung 24:
Einfluss KI-generierter Kampagnenbilder auf die Bewertung des KI-Einsatzes, der Partei
und von KI als Gefahr



Quelle: Eigene Darstellung.

genommene Überzeugungskraft der Bilder und der Bewertung von KI und der Parteien finden. Jedoch zeigt sich hier kein signifikanter Effekt auf die Wahrnehmung des Einsatzes von KI in politischen Kampagnen als Gefahr für die Demokratie.

4.3.6 Einfluss persönlicher Eigenschaften auf die Einstellungen zum Einsatz von KI

Mit Blick auf den Einfluss persönlicher Eigenschaften ist ein zentraler Befund, dass sich sowohl die interne als auch die externe politische Selbstwirksamkeit signifikant positiv auswirkt auf die Bewertung des KI-Einsatzes sowie der Parteien, die KI einsetzen (siehe Abbildung 24). Personen mit hoher politischer Selbstwirksamkeit – also jene, die sich als politisch handlungsfähig wahrnehmen und Vertrauen in das politische System haben – zeigen eine positivere Haltung gegenüber KI und deren Einsatz durch Parteien. Interessanterweise geht eine hohe externe Selbstwirksamkeit mit einer geringeren Wahrnehmung von KI als Gefahr einher. Dies legt nahe, dass das Vertrauen in die Funktionsfähigkeit demokratischer Institutionen potenzielle Bedenken gegenüber dem KI-Einsatz abmildern kann.

Hinsichtlich der demografischen Faktoren zeigt sich ein Altersgradient: Ältere Befragte stehen dem KI-Einsatz kritischer gegenüber, bewerten sowohl die Technologie als auch die sie nutzenden Parteien negativer und nehmen KI eher als Gefahr wahr. Geschlechtsspezifische Unterschiede manifestieren sich in einer signifikant negativeren Bewertung des KI-Einsatzes in Kam-

pagnen durch Frauen im Vergleich zu Männern. Interessanterweise tritt dieser Geschlechtereffekt nicht bei der Parteienbewertung oder der Wahrnehmung von KI als Gefahr auf. Schließlich bewerten Personen mit höherem Kenntnisstand über KI solche Parteien, die KI einsetzen, positiver. Das KI-Wissen steht jedoch nicht im Zusammenhang mit der allgemeinen Bewertung von KI oder deren Wahrnehmung als Gefahr.

4.4 Zwischenfazit: Geringe Erkennung von KI-Bildern und Wirkungsdominanz der Bildinhalte über den KI-Ursprung

Die Ergebnisse unseres Online-Experiments zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern zeichnen ein komplexes Bild. Während der KI-Ursprung von Bildern und entsprechende Kennzeichnungen überraschend wenig Einfluss auf die Wahrnehmung und Bewertung haben, spielen die inhaltliche Gestaltung und die dadurch ausgelösten Emotionen eine zentrale Rolle. Besonders bemerkenswert ist, dass die Wirkung KI-generierter Kampagneninhalte stark von der Übereinstimmung mit den politischen Einstellungen der Rezipient*innen sowie deren persönlichen Eigenschaften wie politischer Selbstwirksamkeit und KI-Kompetenz abhängt. Diese differenzierte Befundlage verdeutlicht, dass eine einfache Unterscheidung zwischen KI-generierten und echten Kampagneninhalten zu kurz greift – vielmehr müssen die komplexen Wechselwirkungen zwischen Bildgestaltung, individuellen Merkmalen und situativen Faktoren berücksichtigt werden.

Im Folgenden fassen wir die zentralen Erkenntnisse entlang unserer Forschungsfragen zusammen.

Fragestellung 7.1: Werden KI-generierte Kampagnenbilder erkannt?

- **Schwierigkeiten bei der Erkennung von KI-Bildern:** Die Teilnehmenden konnten KI-generierte Bilder nur minimal besser von echten Fotografien unterscheiden, was auf den hohen Realitätsgrad von KI-generierten Bildern hinweist.
- **Wirkung von Kennzeichnungen:** Bilder mit einer KI-Kennzeichnung wurden deutlich häufiger als KI-generiert erkannt als Bilder ohne Hinweis, was auf ihre Wirksamkeit hindeutet.
- **Unerwarteter Effekt von Aufklärungstexten:** Ein Aufklärungstext über KI in politischen Kampagnen führte überraschenderweise zu einer geringeren Erkennungsrate von KI-Bildern, was ein Hinweis auf ihre begrenzte Effektivität ist.
- **Einfluss der Bildstimmung:** Negative Bilder wurden häufiger als KI-generiert bewertet als positive, was auf eine mögliche emotionale Wahrnehmungsverzerrung hindeutet.

Fragestellung 7.2: Welche Emotionen lösen KI-generierte Kampagnenbilder aus?

- **Ähnliche emotionale Wirkung von KI-Bildern und echten Fotografien:** KI-generierte Kampagnenbilder lösen fast identische Emotionen wie echte Fotografien aus, mit minimal mehr Wut und Angst sowie etwas weniger Hoffnung.

- **Kein Einfluss durch KI-Disclaimer oder Aufklärungstexte:** Weder optische Hinweise auf KI-Generierung noch Aufklärungstexte beeinflussen die emotionalen Reaktionen der Betrachter*innen.
- **Entscheidende Rolle der Bildinhalte:** Die emotionale Wirkung hängt stark vom Bildinhalt ab. Positive Bilder rufen mehr Hoffnung hervor, während negative Bilder Angst und Wut verstärken.
- **Emotionale Reaktionen auf Migrationsthemen:** Bilder mit Migrationsbezug lösen im Vergleich zu Agrar-Themen weniger Hoffnung, aber mehr Angst und Wut aus.
- **Inhalt vor Herkunft:** Die emotionale Wirkung politischer Kampagnenbilder hängt stärker von ihren Inhalten als von ihrem Ursprung (KI-generiert oder real) ab.

Fragestellung 7.3: Wie werden KI-generierte Kampagnenbilder bewertet?

- **Geringfügige Unterschiede in der Bildbewertung:** KI-generierte Bilder werden als etwas weniger natürlich und überzeugend, jedoch als minimal aufdringlicher empfunden als echte Fotografien.
- **Keine Auswirkungen durch KI-Disclaimer oder Aufklärungstexte:** Weder optische Hinweise noch Aufklärungstexte beeinflussen die Bewertung der Überzeugungskraft oder Authentizität. Hinweislabels führen lediglich zu einer minimal geringeren Wahrnehmung von Natürlichkeit.
- **Höhere Bedeutung von Bildinhalten:** Die Valenz und das Thema der Bildinhalte beein-

flussen die Wahrnehmung stärker als die Entstehungsweise der Bilder. Positive Bilder werden als unaufdringlicher, überzeugender und authentischer wahrgenommen als negative.

- **Wahrnehmung von Themen:** Bilder zu Migrationsthemen werden als aufdringlicher und weniger überzeugend empfunden als solche zu Agrar-Themen, was mit visuellen Framing-Effekten erklärt werden kann (Geise & Lobinger, 2015; Iyer et al., 2014; Nabi, 2003).
- **Übertragene Emotionen auf Bildbewertung:** Positiv oder negativ konnotierte Themen beeinflussen die Gefühlslage der Betrachtenden, was sich direkt auf die Bewertung von Authentizität, Überzeugungskraft und Unaufdringlichkeit der Bilder auswirkt.

Forschungsfrage 8.1: Wie beeinflusst die Rezeption KI-generierter Bilder die Bewertung des Einsatzes von KI in Kampagnen?

- **Einfluss der wahrgenommenen Glaubwürdigkeit:** Befragte, die KI-Kampagnenbotschaften als überzeugend und natürlich wahrnahmen, bewerteten den KI-Einsatz in politischen Kampagnen positiver. Glaubhafte und authentische KI-Inhalte können Skepsis mindern und die Akzeptanz steigern.
- **Rolle der emotionalen Reaktionen:** Positive Emotionen wie Hoffnung führen zu einer positiveren Einschätzung des KI-Einsatzes, während negative Emotionen wie Wut oder Angst die Bewertung verschlechtern.
- **Themen- und Einstellungsdiskrepanz:** Eine Abweichung zwischen den in den KI-Bildern

dargestellten Themen und den politischen Einstellungen der Befragten führt zu kritischen Haltungen gegenüber KI in Kampagnen. Inhalte, die nicht mit der eigenen Weltsicht übereinstimmen, fördern Misstrauen.

- **Einfluss persönlicher Eigenschaften:** Befragte mit hoher interner und externer politischer Selbstwirksamkeit stehen dem KI-Einsatz positiver gegenüber, was in Verbindung mit einem höheren Vertrauen in die eigenen Fähigkeiten und das politische System in Verbindung stehen kann.
- **Demografische Unterschiede:** Ältere Teilnehmende und Frauen bewerteten den KI-Einsatz in politischen Kampagnen tendenziell negativer, was möglicherweise auf geringere technologische Vorkenntnisse und eine damit einhergehende Skepsis zurückzuführen ist.

Forschungsfrage 8.2: Wie beeinflusst die Rezeption KI-generierter Bilder die Bewertung von Parteien, die KI einsetzen?

- **Kein Einfluss durch KI-Ursprung und Hinweise:** Der Ursprung der Bilder (KI-generiert oder nicht) sowie entsprechende Hinweise beeinflussen die Bewertung von Parteien mit KI-Nutzung nicht.
- **Positiver Effekt durch natürliche und überzeugende Inhalte:** Als authentisch und realitätsnah wahrgenommene KI-Bilder steigern die positive Bewertung von Parteien, die KI einsetzen. Überzeugende Inhalte tragen ebenfalls zu einer positiven Bewertung dieser Parteien bei.

- **Emotionale Reaktionen beeinflussen Parteienbewertung:** Die empfundene Hoffnung durch die Rezeption von KI-Bildern führt zu einer positiveren Einstellung gegenüber Parteien, während Wut und Angst eine negative Wirkung haben.
- **Negative Wirkung der positiven Bildvalenz:** Während die Valenz der gesehenen KI-generierten Bilder keinen Einfluss auf die Bewertung des KI-Einsatzes und Gefahrenereinschätzung hat, führen vermehrt gesehene positive KI-Inhalte zu einer kritischeren Bewertung gegenüber Parteien, die KI einsetzen.
- **Wissen und Selbstwirksamkeit als positive Faktoren:** Höheres Wissen über KI und eine erhöhte politische Selbstwirksamkeit (intern und extern) fördern eine positive Bewertung von Parteien, die KI einsetzen.
- **Demografische Unterschiede:** Ältere Befragte bewerten Parteien, die KI nutzen, tendenziell kritischer, was auf Skepsis oder mangelnde Vertrautheit mit neuen Technologien zurückzuführen sein könnte.

Forschungsfrage 8.3: Wie beeinflusst die Rezeption KI-generierter Bilder die Bewertung von KI in Kampagnen als Gefahr?

- **Kein Einfluss durch KI-Ursprung oder Kennzeichnungen:** Weder der Ursprung noch die

Kennzeichnung von KI-Bildern beeinflussen die Wahrnehmung von KI in politischen Kampagnen als Gefahr.

- **Emotionale Reaktionen und Gefahrenwahrnehmung:** Empfundene Wut nach der Betrachtung von KI-Bildern erhöht die Einschätzung von KI als Gefahr, während Angst und Hoffnung keine signifikanten Effekte zeigen.
- **Reduktion der Gefahrenwahrnehmung durch positive Wahrnehmung:** Eine eher unproblematische, natürliche oder überzeugende Wahrnehmung von KI-Bildern mindert die Sorge um Risiken im Zusammenhang mit KI in politischen Kampagnen.
- **Einfluss der politischen Einstellung:** Inhalte, die von der eigenen politischen Überzeugung abweichen, verstärken die Wahrnehmung von KI als Gefahr, was mit dem Konzept der kognitiven Dissonanz erklärbar ist (Festinger, 1957).
- **Rolle politischer Selbstwirksamkeit:** Befragte mit hoher interner und externer politischer Selbstwirksamkeit nehmen KI weniger als Gefahr wahr.
- **Demografische Unterschiede:** Ältere Befragte bewerten KI häufiger als größere Gefahr. Dies deutet auf eine kritischere Haltung gegenüber neuen Technologien in dieser Altersgruppe hin.

5 Fazit und Ausblick

Künstliche Intelligenz revolutioniert zunehmend die Art und Weise, wie politische Kommunikation gestaltet und verbreitet wird. Insbesondere die generativen KI-Systeme ChatGPT und Midjourney, die in Sekundenschnelle täuschend echte Texte, Bilder und Videos in großen Massen erstellen können (Gillespie, 2024; Goldstein et al., 2023), werfen fundamentale Fragen für demokratische Willensbildungsprozesse auf. Die weltweite Debatte über ihren Einsatz in politischen Kampagnen (Chowdhury, 2024; Dommett, 2023; Europäisches Parlament, 2024) kreist dabei um drei zentrale Bedenken: die mögliche Überflutung öffentlicher Diskursräume mit KI-generierten Inhalten (Jungherr, 2023; Jungherr & Schroeder, 2023), das Risiko zunehmender gesellschaftlicher Polarisierung durch personalisierte Botschaften (Battista, 2024; Novelli & Sandri, 2024) sowie einen potenziellen Vertrauensverlust in die politische Kommunikation insgesamt (Hameleers et al., 2024; McKay et al., 2024; Vaccari & Chadwick, 2020).

Während auf regulatorischer Ebene erste Schritte unternommen werden – etwa durch den „AI Act“ des Europäischen Parlaments (Europäisches Parlament, 2024) oder freiwillige Fairnessabkommen deutscher Parteien (Spiegel.de, 2024) – fehlt es bislang an empirischen Erkenntnissen darüber, wie Bürger*innen den Einsatz von KI-generierten Kampagnenbotschaften bewerten und

wahrnehmen. Diese Frage gewinnt zusätzlich an Bedeutung, da Social-Media-Plattformen KI zunehmend als festen Bestandteil ihrer Kommunikationsstrategien etablieren (Robertson, 2024; Sato, 2025; Wankhede, 2024) und gleichzeitig die externe Überprüfung von Inhalten reduzieren (Kaplan, 2025).

Das vorliegende Arbeitspapier begegnet diesem Bedarf und zeichnet ein umfassendes Bild der Einschätzung der deutschen Bevölkerung zum Einsatz von generativer KI in politischen Kampagnen. Es präsentiert zudem neue Erkenntnisse über die Wahrnehmung und Wirkung KI-generierter Kampagneninhalte. Die beiden empirischen Untersuchungen basieren dabei auf einer repräsentativen Online-Befragung und einem Online-Experiment, um sowohl allgemeine Einstellungen als auch konkrete Wahrnehmungs- und Wirkungsmechanismen zu untersuchen.

Die repräsentative Befragung hat verdeutlicht, dass die deutsche Bevölkerung bislang kaum mit KI-generierten Inhalten in Berührung gekommen ist, dem Einsatz von KI in politischen Kampagnen gleichwohl überwiegend skeptisch gegenübersteht. Diese skeptische Haltung spiegelt sich einerseits in der negativen Bewertung des generellen Einsatzes von generativer KI und davon ausgehenden Gefahren wider. Andererseits zeigt sie sich gegenüber Parteien, die KI ohne

Kennzeichnung für die Wähler*innenkommunikation und Erstellung von realistischen Bildern sowie Videos einsetzen. Diese Skepsis speist sich vor allem aus Bedenken hinsichtlich möglicher Manipulationen und einem potenziellen Vertrauensverlust in die politische Kommunikation. Unsere Ergebnisse erweitern damit vorhandene Studien, die eine ähnlich kritische Haltung der Bevölkerung gegenüber KI beispielsweise im Journalismus finden. Dies weist auf ein generelles Misstrauen gegenüber KI-Technologien hin, insbesondere wenn KI wichtige Aufgaben übernimmt und in wichtigen Themenbereichen wie der politischen Berichterstattung eingesetzt wird (vgl. Kieslich et al., 2022; Vogler et al., 2023). Darüber hinaus bestätigen unsere Befunde die geäußerten Bedenken im gesellschaftlichen und wissenschaftlichen Diskurs (Chowdhury, 2024; Dommett, 2023; Europäisches Parlament, 2024; Gillespie, 2024; Jungherr, 2023).

Gleichzeitig hatten viele Befragte bisher kaum bewusste Berührungspunkte mit KI-generierten Inhalten in Kampagnen. Diese Diskrepanz zwischen dem tatsächlichen Einsatz von KI durch politische Akteur*innen und der Wahrnehmung in der Bevölkerung wirft Fragen auf. Einerseits kann angenommen werden, dass der Einsatz zum Zeitpunkt der Befragung (Mai 2024) eher sporadisch war und die Befragten eventuell tatsächlich keinen Kontakt zu KI-generierten Kampagneninhalten hatten. Möglich wäre auch, dass die Befragten zwar Kontakt hatten, diesen aber nicht bewusst wahrgenommen haben, weil Inhalte bisher noch nicht gekennzeichnet werden müssen. Entsprechend wäre die mangelnde Kenntnis dann eine Folge fehlender Transparenz.

Bemerkenswert ist zudem, dass trotz fehlender eigener Berührungspunkte mit KI-generierten Inhalten eine hohe Skepsis vorherrscht. Eine mögliche Quelle für eine eher skeptische Sicht könnte die Medienberichterstattung sein, die sich häufig kritisch mit technischen Neuerungen in politischen Kampagnen auseinandersetzt (Jungherr & Rauchfleisch, 2024; MeMo:KI – Medienanalyse, o.J.).

Letztlich zeigen die Befragungsergebnisse aber auch Möglichkeiten zur Nutzung von generativer KI in politischen Kampagnen durch den Wunsch der Bevölkerung nach einem transparenten Umgang mit KI durch die Parteien und einer Regulierung des Einsatzes durch Behörden. So akzeptieren die Befragten den Einsatz von KI am ehesten für sprachliche Überarbeitungs- oder Korrekturaufgaben und insbesondere, wenn es eine Kennzeichnungspflicht für alle KI-generierten Inhalte gibt (für kritische Auseinandersetzung siehe Altay & Gilardi, 2024; Leibowicz, 2024), was im Einklang mit den Anforderungen des europäischen AI Acts (Europäisches Parlament, 2024) und den Selbstverpflichtungen einiger Parteien und Plattformen steht. Dies bestätigt die Befunde aus anderen Studien (Fletcher & Nielsen, 2024; Kieslich et al., 2022; Votta & de León, 2024) und deutet daraufhin, dass die hohe Akzeptanz für unterstützende Funktionen darauf zurückzuführen ist, dass diese Aufgaben als weniger sensibel oder manipulativ in der Bevölkerung wahrgenommen werden. Je stärker jedoch KI in die inhaltliche Gestaltung und Personalisierung der Wählerkommunikation eingreift, desto größer scheinen die Vorbehalte zu sein. Zusammenfassend deuten die Befragungsergebnisse daraufhin, dass wenn

KI nicht transparent und verantwortungsvoll eingesetzt wird, eine große Skepsis bei der Bevölkerung gegenüber dem Einsatz von KI herrschen wird, wodurch letztlich auch ihr Vertrauen in politische Akteur*innen und Institutionen negativ beeinflusst werden kann (Gordon, 2024; Jungherr, 2023).

Mit Blick auf die Wahrnehmung von KI-Kampagnenbildern zeigt das durchgeführte Online-Experiment, dass Schwierigkeiten dabei bestehen, KI-generierte Inhalte von natürlichen Fotografien zu unterscheiden, unabhängig von ihrem KI-Ursprung, der Valenz und eines Aufklärungstextes. Dies steht im Einklang mit früheren Studien (u. a. Frank et al., 2023; Groh et al., 2024; Köbis et al., 2021), die ebenfalls große Schwierigkeiten bei der Unterscheidung zwischen KI-generierten und echten Inhalten festgestellt haben. Gleichwohl zeigt unsere Studie, dass klare KI-Kennzeichnungen im Bild die Erkennung unterstützen und bestätigt damit die Befunde von Lu und Yuan (2024), die ebenfalls einen positiven Einfluss von KI-Disclaimern feststellten. Überraschenderweise verringern Aufklärungstexte die Wahrscheinlichkeit, Bilder als KI-generiert zu erkennen. Dies kann einerseits daran liegen, dass unser Aufklärungstext zu viele oder komplizierte Informationen beinhaltet hat oder er schlicht von Befragten übersehen beziehungsweise nicht gelesen wurde. Andererseits deuten auch Ergebnisse anderer Studien (Groh et al., 2024; Kasra et al., 2018; Köbis et al., 2021) an, dass Vorwissen oder Aufklärung nicht unbedingt zu einer besseren Erkennung führt (siehe aber Tahir et al., 2021). Gleichwohl bedarf es weiterer Forschung, die sich mit den Rahmenbedingungen der Wirksam-

keit von Hinweisen und Kennzeichnungen befassen. Denn Disclaimer dürfen nicht als universelles „Allheilmittel“ betrachtet werden, sondern müssen mit weiteren Aufklärungsmaßnahmen einhergehen (für kritische Auseinandersetzung siehe Altay & Gilardi, 2024; Leibowicz, 2024).

Unser Experiment zeigt darüber hinaus, dass sich KI-generierte Kampagnenbilder in ihrer emotionalen Wirkung und Überzeugungskraft kaum von echten Bildern unterscheiden. Dies steht im Einklang mit früheren Forschungen, die ähnliche emotionale Reaktionen auf KI-generierte und echte Bilder feststellten (Eiserbeck et al., 2023; Nightingale & Farid, 2022). Gleichwohl sollte bei diesen Befunden beachtet werden, dass wir bei der Erstellung der Bilder bedacht waren, die Darstellung vergleichbar mit realen Bildern zu halten. Eine stärkere emotionale Aufladung ist durch entsprechendes Prompting problemlos möglich. Die Ergebnisse unserer Untersuchung verweisen dabei auf einen deutlichen Effekt der Valenz des Inhalts auf Emotionen und die Wahrnehmung der Inhalte. Trotz der vergleichsweise milden Manipulation der Bilder lösten positive Anzeigen mehr Hoffnung und negative Bilder mehr Angst sowie Wut aus. Dies knüpft an frühere Forschungsergebnisse zur emotionalen Wirkung von Bildern in der politischen Kommunikation an (Geise, 2011; Leonhard & Bartsch, 2020). Außerdem wurden positive Anzeigen im Vergleich zu negativen als unaufdringlicher, überzeugender und authentischer wahrgenommen. Bei dieser Wahrnehmung spielte es ebenfalls keine Rolle, ob es sich um ein „echtes“ oder KI-generiertes Bild handelte. Der negative Effekt ist auch durch Forschung zur Wirkung visueller Darstellungen

in der „klassischen“ politischen Kommunikation bekannt (Geise, 2011; Maurer, 2016). Damit unterstützen die Ergebnisse die Annahme, dass KI-Technologien die Grenzen zwischen authentischer und künstlicher Kommunikation verwischen können (Eiserbeck et al., 2023; Nightingale & Farid, 2022), zeigen aber gleichzeitig, dass „klassische“ Faktoren der politischen Kommunikation keineswegs an Gültigkeit verlieren.

Neben der Bewertung und emotionalen Wirkung der Inhalte war ein weiteres Anliegen zu klären, ob die Konfrontation mit teils KI-generierten politischen Anzeigen die generelle Wahrnehmung von KI beeinflusst. Dazu haben wir aus den jeweiligen Reaktionen auf alle drei gezeigten Inhalte gemittelt und den Einfluss dieser durchschnittlichen Einschätzungen auf die Wahrnehmung von KI und die Bewertung von Parteien, die KI einsetzen, gemessen. Dabei zeigt sich, dass emotionale Reaktionen eine zentrale Rolle bei der Akzeptanz von KI spielen (Beierlein & Burger, 2019b). Je mehr Wut und Angst die Teilnehmenden empfunden haben, desto mehr wurde KI als Bedrohung wahrgenommen, was durch die Forschung von Eiserbeck et al. (2023) bestätigt wird, die ebenfalls zeigt, dass negative Darstellungen die Risikowahrnehmung erhöhen können. Gleichzeitig sinkt auch die Bewertung von Parteien: Menschen, die überwiegend negative Emotionen bei der Rezeption empfunden haben, haben ein negativeres Bild von Parteien, die KI einsetzen und sehen im Einsatz von KI eher Gefahren. Eine übermäßige Emotionalisierung – beispielsweise, um Panik oder Wut auf den politischen Gegner zu erzeugen – könnte so zum Bumerang für die Parteien werden. Neben den Emotionen spielt

auch die Wahrnehmung der Authentizität eine zentrale Rolle. Wurden die Bilder als eher unproblematisch wahrgenommen, verbesserte sich die Bewertung der Parteien, die KI einsetzen, während problematisch wahrgenommene Bilder die Akzeptanz reduzierten. Diese Ergebnisse decken sich mit den Befunden von Dobber et al. (2021), die zeigen, dass die Authentizität von KI-generierten Inhalten entscheidend für die Glaubwürdigkeit und Akzeptanz ist.

Schließlich offenbart unsere Studie einen Einfluss der Übereinstimmung zwischen den dargestellten Themen in KI-generierten Bildern und den politischen Einstellungen der Betrachter*innen. Es zeigt sich, dass eine größere Diskrepanz zwischen dem präsentierten Thema und den eigenen politischen Überzeugungen zu einer signifikant kritischeren Haltung gegenüber dem KI-Einsatz führt und gleichzeitig die Wahrnehmung von KI als potenzielle Gefahr verstärkt. Die Abwertung von dissonanten Inhalten deckt sich mit Forschungsergebnissen zum „Motivated Reasoning“ (Kunda, 1990). Für politische Akteur*innen eröffnet KI die Möglichkeit, Kampagneninhalte besonders effektiv zu gestalten, indem Motive und Bildsprache so gewählt werden können, dass diese mit den politischen Einstellungen spezifischer Zielgruppen übereinstimmen. Dabei besteht jedoch die Gefahr, dass KI-generierte Desinformationen unkritisch rezipiert beziehungsweise nicht als solche erkannt und weiterverbreitet werden, solange die transportierte Botschaft das eigene Weltbild stützt (Duffy et al., 2020; Sindermann et al., 2020). Das könnte Polarisierungstendenzen verstärken (Hameleers & Van der Meer, 2020). Dies ist mitnichten ein neu-

es Phänomen; dennoch birgt generative KI durch die einfache Nutzbarkeit, Kosteneffizienz und Schnelligkeit der Erzeugung maßgeschneiderter und kaum von realen Inhalten unterscheidbarer Inhalte großes Potenzial – eben auch für Missbrauch (Battista, 2024; Jungherr & Schroeder, 2023; Novelli & Sandri, 2024).

Die Untersuchung bemühte ein aufwendiges Design und achtete bei der Erstellung der Stimuli auf eine realistische Anmutung. Gleichwohl gibt es Limitationen, die bei der Interpretation der Ergebnisse zu berücksichtigen sind und in zukünftigen Forschungsarbeiten adressiert werden könnten: Im Hinblick auf unsere Befragungsstudie ist zunächst die Frage der Repräsentativität zu diskutieren. Obwohl wir uns um eine quotierte Stichprobe bemühten, können wir Verzerrungen durch Selbstselektion und die ausschließlich online durchgeführte Rekrutierung nicht vollständig ausschließen. Eine Kombination aus Online- und Offline-Befragungsmethoden könnte in Zukunft zu einer verbesserten Repräsentativität beitragen und ein breiteres Spektrum der Bevölkerung einbeziehen. Zudem mussten unsere Befragten teilweise Einschätzungen zu hypothetischen KI-Einsatzszenarien abgeben. Die Validität solcher Antworten könnte durch mangelnde persönliche Erfahrung oder eingeschränkte Vorstellungskraft der Teilnehmenden beeinträchtigt sein. Dies wird sich durch den zunehmenden Einsatz von KI – sofern er denn als solcher identifiziert wird – vermutlich ändern.

Unser Experiment konzentrierte sich auf die Messung kurzfristiger Effekte. Wir haben lediglich unmittelbare Reaktionen erfasst, während

kumulative Effekte mehrmaliger Exposition und mögliche Einstellungsänderungen über einen längeren Zeitraum nicht berücksichtigt wurden. Längsschnittstudien könnten hier zusätzliche Erkenntnisse liefern. Eine weitere Einschränkung unseres Experiments liegt in der Beschränkung auf visuelle Reize, insbesondere statische Bilder. Angesichts der zunehmenden Bedeutung von Videos, Audioformaten und interaktiven Inhalten in modernen politischen Kampagnen stellt dies eine klare Limitation dar. Zukünftige Studien sollten ein breiteres Spektrum an KI-generierten Medienformaten einbeziehen, um der Komplexität der politischen Kommunikationslandschaft gerecht zu werden. Zudem beziehen sich die Resultate auf spezifische, von uns erstellte KI-Bilder. Wenn gleich wir insgesamt über 400 verschiedene Varianten von Anzeigen erstellt haben, bleibt offen, inwieweit sie tatsächlich auf andere KI-generierte Inhalte oder reale Kampagnenmaterialien übertragbar sind. Dies bedarf weiterer Untersuchungen. Überdies haben wir auf eine übertriebene Emotionalisierung verzichtet. Gerade das Potenzial von KI, erfundene Tatsachen auf drastisch emotionalisierte Weise darzustellen, birgt Gefahren, die es zukünftig zu untersuchen gilt. Schließlich haben wir nicht nur die Bilder, sondern auch die Texte entsprechend der Valenz des Bildes angepasst. Der durchgehend starke Effekt der Valenz kann entsprechend nicht nur der Bildmanipulation zugeschrieben werden. Anknüpfende Forschung könnte sich auf die ausschließliche Manipulation der Bilder beschränken, was gleichwohl Fragen der externen Validität aufwerfen könnte.

Trotz dieser Einschränkungen zeigen die Ergebnisse der Befragung, dass die Risiken des KI-Ein-

satzes in der politischen Kommunikation von der Bevölkerung wahrgenommen und entsprechende Regulierungen befürwortet werden, wenngleich die Kenntnisse zu KI bei einigen Menschen nur bedingt vorhanden sind. Im Experiment wird gleichzeitig deutlich, dass die Menschen KI-generierte Bilder kaum von „echten“ Bildern unterscheiden können; gleichwohl hilft die Kennzeichnung der Bilder bei der Erkennung. Diese und andere Erkenntnisse haben Implikationen für unterschiedliche Stakeholder*innen in demokratischen Gesellschaften (siehe Kapitel 6).

Insgesamt deuten die Ergebnisse darauf hin, dass KI eher als Verstärker bestehender Kommunikationsstrategien wirkt, als dass sie grundlegend neue Wirkungs dynamiken schafft. Jene, die die Technologie für undemokratische Ziele nutzen wollen, bekommen mit KI neue Möglichkeiten der Manipulation an die Hand. Aber auch schon vor dem Einsatz von KI kommunizierten einige Akteur*innen emotional und verbreiteten Halb- oder Unwahrheiten. Auch wurde früher auf Wahlplakaten mit bezahlten Modellen und nicht mit beispielsweise echten Handwerker*innen geworben. Das legt nahe, dass die Hauptherausforderungen im Zusammenhang mit KI in politischen Kampagnen weniger in der Technologie selbst liegen, sondern vielmehr in der Art und Weise, wie sie eingesetzt wird. Bestehende Tendenzen wie die Personalisierung von Botschaften, emotionale Appelle oder die gezielte Ansprache spezifischer Wählergruppen könnten durch KI effizienter und in größerem Maßstab umgesetzt werden. Dies könnte

einerseits die politische Kommunikation effektiver machen, andererseits aber auch Risiken wie eine verstärkte Polarisierung oder die Verbreitung von Desinformation begünstigen. Die Ergebnisse unterstreichen, dass die Herausforderungen durch KI in der politischen Kommunikation nicht primär technischer, sondern vor allem gesellschaftlicher und ethischer Natur sind. Sie verdeutlichen die Notwendigkeit eines gesamtgesellschaftlichen Dialogs über den verantwortungsvollen Umgang mit KI in demokratischen Prozessen.

Um diesen Dialog informiert führen zu können, ist weitere Forschung nötig. Dabei sollte in den Blick genommen werden, wie sich Einstellungen und Kompetenzen im Umgang mit KI-generierten Inhalten über Zeit entwickeln und welche langfristigen Auswirkungen der KI-Einsatz auf das Vertrauen in politische Institutionen und Prozesse haben kann. Schließlich gilt es abzuwägen und eine Balance zwischen Innovation und dem Schutz demokratischer Werte zu finden. Wer die Demokratie vor dem Missbrauch von KI schützen will, muss sich für langfristige Schutzmaßnahmen einsetzen. Dazu gehören eine umfangreiche gesellschaftliche Aufklärung und Förderung der KI- und Medienkompetenz, ebenso wie effektive Ansätze zur Regulierung des Einsatzes von KI in der politischen Kommunikation und nicht zuletzt der Appell zu einem verantwortlichen Umgang seitens politischer Akteur*innen. Einige Vorschläge für einen verantwortungsbewussten Umgang mit KI in der politischen Kommunikation formulieren wir im abschließenden sechsten Kapitel.

6 Vorschläge für einen verantwortungsbewussten Umgang mit KI in politischen Kampagnen

Die Ergebnisse unserer Studien haben gezeigt, dass der Einsatz von generativer KI in politischen Kampagnen sowohl Chancen als auch Risiken für Bürger*innen, politische Akteur*innen und demokratische Prozesse birgt. Um das Potenzial dieser Technologie verantwortungsvoll zu nutzen und gleichzeitig mögliche Gefahren zu minimieren, sind Handlungsvorschläge erforderlich. Auf Grundlage unserer empirischen Erkenntnisse zur Akzeptanz, Wahrnehmung und Wirkung von KI in der politischen Kommunikation skizzieren wir im Folgenden Handlungs-ideen, die Bürger*innen, zivilgesellschaftliche Akteur*innen, politische Akteur*innen, Journalist*innen sowie Regulierungsakteur*innen beim verantwortungsbewussten Umgang mit KI in politischen Kampagnen unterstützen können.

6.1 Bürger*innen und zivilgesellschaftliche Akteur*innen: Förderung der Medienkompetenz im Umgang mit KI in politischen Kampagnen

Bürger*innen müssen sich auf einen vermehrten Einsatz generativer KI in politischen Kampagnen einstellen. Sie stehen vor der Herausforderung, KI-generierte Inhalte zu erkennen und damit kritisch umzugehen. Für die Bewältigung dieser Aufgabe können zivilgesellschaftliche Akteur*innen eine wichtige Rolle einnehmen.

KI-Kompetenz und -Wissen fördern

Bürger*innen können selbst aktiv daran arbeiten, ihre Medienkompetenz mit Blick auf KI zu erweitern oder zu verbessern. Neben der Medienkunde, -kritik, -nutzung und auch -gestaltung (Bateman & Jackson, 2024; Hugger, 2022; Iske & Barberi, 2022), sollten die Funktionsweisen, Einsatzfelder und Nutzungsmöglichkeiten von KI wesentlicher Bestandteil der Medienkompetenz sein (Hwang et al., 2021; Wang et al., 2024). Denn unsere Forschungsergebnisse zeigen: das Wissen von Bürger*innen über die Einsatzmöglichkeiten von KI in politischen Kampagnen und die Fähigkeit zur Erkennung von KI-generierten Inhalten ist in einigen Bereichen unzureichend.

Zivilgesellschaftliche Akteur*innen spielen eine zentrale Rolle bei der Vermittlung dieser Kompetenzen. Sie können niedrigschwellige und (medien)pädagogische Bildungsangebote schaffen, beispielsweise interaktive (Online-) Kurse erstellen, (Online-)Workshops organisieren und als Multiplikatoren fungieren. Darüber hinaus können sie Fact-Checking-Tools entwickeln, Plattformen zur Erkennung von KI-Inhalten bereitstellen und Bürger*innenlabore einrichten, in denen Menschen KI-Technologien praktisch erproben können. Auch die Entwicklung von „Serious Games“ zur spielerischen Vermittlung von KI-Kompetenzen, die Organisation von Hackathons für demokratiefördernd-

de KI-Anwendungen sowie die Einrichtung von lokalen Tech-Cafés gehören zu möglichen Formaten. Besonders wichtig ist ihre Brückenfunktion zwischen technologischer Entwicklung und gesellschaftlicher Vermittlung. Denn nur wenn die Bildungsangebote im Gleichschritt mit den technologischen Entwicklungen gehen, können die komplexen und sich schnell entwickelnden KI-Themen verständlich aufbereitet und in die Breite der Gesellschaft getragen werden.

Gute, oft kostenlose Bildungsangebote zum Thema Medienkompetenz lassen sich schon aktuell bei Landesmedienanstalten, der Bundes- beziehungsweise den Landeszentralen für politische Bildung, Bundes- und Landessämtern für Informationssicherheit oder Datenschutz sowie anderen Bildungseinrichtungen (Hochschulen, Erwachsenenbildungsträgern, Museen, Bildungswerken) finden. Darüber hinaus gibt es auch spielerische Internetformate zur Förderung der Medienkompetenz mit einem KI-Schwerpunkt. Beispielsweise hier:

- Escape Game der KAS: https://bit.ly/OBS_GenAI1
- Detect AI Fakes Game der Northwestern University: https://bit.ly/OBS_GenAI2
- Newstest der Medienanstalt Berlin-Brandenburg: https://bit.ly/OBS_GenAI3
- Digital Check der Gesellschaft für Medienpädagogik und Kommunikationskultur: https://bit.ly/OBS_GenAI4

- Fit for News des Europäische Instituts für Journalismus- und Kommunikationsforschung: https://bit.ly/OBS_GenAI5

Kritischer und reflektierter Medienkonsum

Unsere Untersuchung hat gezeigt, dass politische Kampagnen vermehrt KI-generierte Inhalte in ihrer Kommunikation einsetzen (siehe 2.1). Daher ist es wichtiger denn je, dass sich Bürger*innen mit politischen Informationen, die sie auf Plakaten oder über Massenmedien beziehen, noch intensiver und kritischer auseinandersetzen. Unsere Studie zeigt, dass insbesondere Bildinhalte wie die Valenz oder Themen die Botschaftswahrnehmung und die Einstellungen gegenüber KI in politischen Kampagnen beeinflussen. Daher sollten politische Informationen, insbesondere in Wahlkampfzeiten oder während Kriegen und Anschlägen, mit besonderer Sorgfalt überprüft werden. Neben passenden Angeboten von zivilgesellschaftlichen Akteur*innen, gibt es Möglichkeiten zur eigenständigen Überprüfung der Echtheit von Bildern und Videos, beispielsweise mit einer Suche über Googles Fact-Check Explorer, der Überprüfung mit Informationen aus verschiedenen etablierten Medien und Fact-Checking-Plattformen oder mithilfe von KI-Tools zur Erkennung von Deepfakes (z. B. [forensics.com](https://www.forensics.com); [deepware.ai](https://www.deepware.ai); [sightengine.com](https://www.sightengine.com)).⁶ Darüber hinaus sollten Bilder und Videos auf Anzeichen für eine KI-Generierung überprüft werden (siehe Infobox 5), beispielsweise anatomische Fehler, unnatürlich glatte oder faltige Haut, feh-

6 Wichtig ist hierbei anzumerken, dass die KI-Tools zur automatischen Erkennung von Deepfakes fehleranfällig sind und keine absolute Sicherheit zum KI-Ursprung geben. Sie sollten daher nur als Anhaltspunkt für einen möglichen KI-Ursprung genutzt werden.



Infobox 5:

Anzeichen für ein KI-generiertes Bild (nach Kamali et al., 2024)

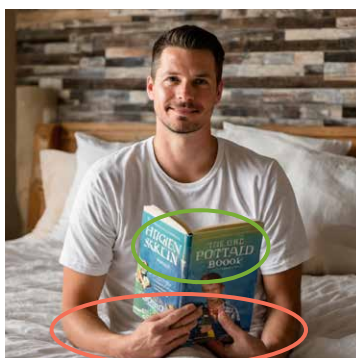
Sozialkulturelle Ungereimtheiten

(bspw. unwahrscheinliche Szenarien, nicht zur abgebildeten Person passende Verhaltensweisen)



Physische Ungereimtheiten

(bspw. fehlerhafte Spiegelbilder, Proportionen und Perspektive von Bildinhalten, Schwerkraft)



Anatomische Ungereimtheiten

(bspw. fehlende Körperteile, zu viele Körperteile, verschwommene Gesichter, nicht symmetrische Gesichtspartien)

Funktionale Ungereimtheiten

(bspw. Fantasieschriften; Fantasiemarken; Fantasievereine)

lerhafte Schattenwürfe, unrealistische Reflexionen auf Brillen, auffällige Unstimmigkeiten bei Blinzeln und Lippenbewegungen beim Sprechen (für eine ausführliche Übersicht siehe Kamali et al., 2024). Zusätzlich können Bürger*innen auf der Plattform campaigntracker.de nachverfolgen, welche KI-generierten Social-Media-Botschaften von deutschen politischen Akteur*innen im Umlauf sind (Kruschinski et al., 2025).

Aktiv sein gegen schädliche KI-Nutzung

Bürger*innen sollten die in der Umfrage deutlich gewordenen Bedenken äußern und von politischen Akteur*innen und Technologieunternehmen mehr Transparenz und Rechenschaftspflicht im Umgang mit KI in politischen Prozessen einfordern. Sie sollten eine aktive Rolle im Kampf gegen KI-generierte Desinformation übernehmen, indem sie verdächtige Inhalte an Plattformen melden und potenzielle Verstöße den zuständigen Wahl- oder Datenschutzbehörden mitteilen.

Zivilgesellschaftliche Akteur*innen können diese individuellen Bemühungen durch kollektives Handeln verstärken: Sie können Watchdog-Funktionen übernehmen, systematische Monitoring-Systeme für KI-generierte Desinformation entwickeln und entsprechende Meldestellen einrichten. Durch die Bildung von Netzwerken und Bündnissen können sie den Druck auf politische Entscheidungsträger*innen und Tech-Unternehmen erhöhen. Besonders wichtig ist auch ihre Rolle bei der rechtlichen Unterstützung von Bürger*innen, etwa durch Musterklagen gegen missbräuchliche KI-Nutzung oder durch Beratung zur Durchsetzung von Datenschutzrechten.

Außerdem sollten Bürger*innen auch verstärkt ihr Recht auf Einsicht in die Nutzung ihrer Daten und ihres Widerspruchs dagegen Gebrauch machen (<https://www.datenschutz-berlin.de/buergerinnen-und-buerger/selbstdatenschutz/>). Zivilgesellschaftliche Datenschutzorganisationen können hier durch Aufklärungskampagnen, die Bereitstellung von Musterformularen und konkrete Hilfestellung bei der Wahrnehmung dieser Rechte unterstützen.

6.2 Politische Akteur*innen: KI verantwortungsvoll in Kampagnen einsetzen

Politische Akteur*innen spielen eine Schlüsselrolle bei der Nutzung von KI in der politischen Kommunikation. Sie stehen vor der Herausforderung, innerhalb des bestehenden rechtlichen Rahmens KI effektiv einzusetzen und dabei gleichzeitig verantwortungsvoll mit den Möglichkeiten und Risiken dieser Technologie umzugehen.

Transparenzmaßnahmen in KI-Richtlinien verankern

Die repräsentative Befragung zeigt, dass Bürger*innen dem Einsatz von KI in politischen Kampagnen skeptisch bis ablehnend gegenüberstehen. Sie würden einen Einsatz jedoch am ehesten akzeptabel finden, wenn KI-generierte Texte, Bilder und Videos sichtbare KI-Kennzeichnungen beinhalten. Solche Transparenzhinweise zeigten in unserem Online-Experiment auch keine signifikanten Einflüsse auf die Bewertung von Parteien, die KI einsetzen oder auf die Einschätzung des Einsatzes von KI als Gefahr für

die Demokratie. Entsprechend sollten politische Akteur*innen eindeutig erkennbare Kennzeichnungen auf ihren KI-Inhalten verwenden, um Menschen auf die KI-Herkunft hinzuweisen. Denn die negativen Folgen einer aufgedeckten, nicht gekennzeichneten KI-Nutzung könnten möglicherweise schwerwiegender für die Bewertung der Partei sein als ein transparenter Umgang von Anfang an (siehe Weikmann et al., 2024). Diese Kennzeichnungen ermöglichen es den Bürger*innen, die Inhalte kritisch zu reflektieren, ohne einen Einfluss auf die Parteienbewertung zu nehmen. Wie bereits bei einigen Parteien geschehen, sollten diese Kennzeichnungsmaßnahmen in Richtlinien zur Nutzung von generativer KI in der Partei- und Kampagnenarbeit festgeschrieben werden (siehe Kapitel 2.4). Dies trägt zur Stärkung des Vertrauens in den verantwortungsvollen Umgang mit KI sowie der Sicherung der Glaubwürdigkeit politischer Kommunikation bei.

Ethische Standards für KI-Nutzung festlegen

Insbesondere bei den KI-Inhalten, die die Kommunikation mit Wähler*innen betreffen, sollten politische Akteur*innen darauf achten, ethische Standards einzuhalten. Dieser Vorschlag basiert zum einen auf dem Ergebnis, dass der Einsatz von KI – insbesondere für Bilder und Videos – von einer Mehrheit der Befragten kritisch gesehen wird. Zum anderen zeigt unser Online-Experiment, dass KI-Botschaften mit negativem Inhalt zu negativeren Emotionen und Einstellungen gegenüber dem KI-Einsatz führt. Daher sollten politische Akteur*innen sicherstellen, dass die generierten Inhalte den Grundsätzen der Fair-

ness und ethischen Unbedenklichkeit entsprechen und nicht dazu genutzt werden, gezielt emotionale Schwächen auszunutzen oder Desinformationen über politische Gegner*innen zu verbreiten. Hier wäre eine Selbstverpflichtung zu ethischen Leitlinien wünschenswert, die solche Praktiken ausschließt, um die Integrität der politischen Kommunikation trotz der Nutzung von KI zu wahren (Chowdhury, 2024; G'sell, 2024; Ryan, 2020; Taddeo & Floridi, 2018).

Einstellungen der Bürger*innen zu KI berücksichtigen

Politische Akteur*innen sollten die Risiken des Einsatzes von KI in der politischen Kommunikation durch den Einsatz von Monitoring- und Feedback-Mechanismen kontinuierlich bewerten. Insbesondere die repräsentative Befragung hat gezeigt, dass Bürger*innen KI in der politischen Kommunikation oft kritisch sehen und teils sogar ängstlich darauf reagieren. Politische Akteur*innen sollten daher regelmäßig Rückmeldungen von Bürger*innen einholen, um ihre Strategien anzupassen und negative Wahrnehmungen frühzeitig zu erkennen und anzugehen. Die Implementierung von Feedback-Schleifen und die offene Kommunikation über Anpassungen aufgrund von Bürger*innenfeedback können dazu beitragen, dass die Nutzung von KI nicht als intransparente oder manipulative Technik wahrgenommen wird. Stattdessen wird die Bereitschaft gezeigt, auf die Bedenken der Wählerschaft einzugehen und den Einsatz von KI in der politischen Kommunikation stetig zu verbessern. Diese proaktive Haltung kann maßgeblich dazu beitragen, Vertrauen zu schaffen und die Akzeptanz des KI-Einsatzes in

politischen Kampagnen zu erhöhen (Bateman & Jackson, 2024; McKay et al., 2024).

6.3 Journalist*innen: Über KI-Einsatz mit Wissen und Evidenz berichten

Journalist*innen nehmen eine zentrale Funktion bei der öffentlichen Vermittlung und der Einordnung des Einsatzes von generativer KI in der politischen Kampagnenarbeit ein. Sie stehen vor der Aufgabe, komplexe technologische Entwicklungen verständlich zu machen und dabei sowohl kritisch als auch ausgewogen zu berichten. Zwei spezifische Aufgaben für den Journalismus werden durch unsere Ergebnisse unterstrichen:

Evidenzbasierte Berichterstattung ohne Alarmismus zu Gefahren von KI

Unsere Untersuchungen zeigen, dass KI-generierte Inhalte zwar potenziell schädlich eingesetzt werden können, die tatsächlichen Auswirkungen – so zeigt unser Online-Experiment – aber von vielfältigen Faktoren wie der Kommunikationsstrategie, dem*der Empfänger*in oder dem Kontext der Rezeption abhängen. Daher sollten Journalist*innen nur dann über den Einsatz von KI für Desinformation oder Negative Campaigning berichten, wenn dafür klare Belege vorliegen. Aktuelle Studien belegen nämlich, dass die Diskussion über den Kampagneneinsatz von KI und anderen modernen Technologien in der Medienberichterstattung eher negativ geführt wird (Sanguinetti & Palomo, 2024; Simon, Altay, & Mercier, 2024). Expert*innen

argumentieren, dass eine alarmistische Berichterstattung über KI-generierte Deepfakes und die damit verbundene Warnung vor einer ‚Informationsapokalypse‘ zu einer sich selbst erfüllenden Prophezeiung werden könnte. Sie regen daher einen sachlicheren und evidenzbasierten journalistischen Diskurs an (Tabuz & Nehring, 2024). Darüber hinaus sollten die Einflüsse des Einsatzes von KI auf politisches Verhalten oder demokratische Prozesse mit wissenschaftlicher Evidenz und ohne unbegründeten Alarmismus berichtet werden. Übertreibungen in den Medienberichten könnten das bereits bestehende Misstrauen gegenüber KI, ihrem Einsatz in der Politik und demokratischen Prozessen weiter verstärken (Jungherr & Rauchfleisch, 2024). Außerdem zeigt unser Online-Experiment, dass Wissen über KI mit einer positiveren Bewertung von Parteien, die KI einsetzen, einhergeht. Journalist*innen könnten auch über die positiven Einsatzmöglichkeiten von KI berichten, etwa wie diese Technologien dazu beitragen können, Wähler*innen besser zu informieren und politische Prozesse transparenter zu gestalten (Bieber et al., 2024).

KI-Kompetenzen von Journalist*innen aufbauen und stärken

Medienorganisationen sollten die Kompetenzen und Werkzeuge ihrer Journalist*innen im Bereich KI gezielt fördern, damit sie eine potenziell schädliche Nutzung generativer KI erkennen und verständlich darüber berichten können. Unsere Befragungsergebnisse zeigen, dass Menschen bisher nur in geringem Maße

mit KI in Berührung gekommen sind, eine durchaus solide KI-Kompetenz besitzen, aber nur begrenzt KI-Inhalte erkennen können. Daher ist es von großer Bedeutung, dass es Journalist*innen gelingt, Einsatz und Wirkung generativer KI für die Öffentlichkeit zu vermitteln und Chancen sowie Gefahren einzuordnen (Bateman & Jackson, 2024). Dabei stehen sie vor der Herausforderung, komplexe technologische Entwicklungen für die Allgemeinheit verständlich zu machen. Es ist folglich wichtig, dass Journalist*innen technisches Wissen und Werkzeuge besitzen, um eine schädliche Nutzung nachzuweisen, Folgen realistisch abzuwägen und Verantwortliche dafür benennen zu können (Łabuz & Nehring, 2024). Dies kann nicht nur dazu beitragen, die Öffentlichkeit über KI-Kampagnen zu informieren, sondern auch die Integrität des Informationsökosystems zu schützen sowie das Vertrauen in demokratische Prozesse und Institutionen zu stärken (Jungherr & Rauchfleisch, 2024; Vaccari & Chadwick, 2020).

6.4 Regulierungsakteur*innen: Evidenzbasierte Richtlinien für den Einsatz von KI in politischen Kampagnen schaffen und durchsetzen

Regulierungsakteur*innen nehmen eine Schlüsselrolle für den Einsatz von KI in politischen Kampagnen ein. Sie stehen vor der schwierigen Aufgabe, einen rechtlichen Rahmen für den Einsatz von KI zu schaffen (u.a. Bateman & Jackson, 2024; G'sell, 2024; Jungherr, 2024; Laude & Daum, 2024; Taddeo & Floridi, 2018).

Wirkungsvolle Regulierungsrahmen für den Einsatz von KI in Wahlprozessen schaffen

Unsere Ergebnisse zeigen, dass eine klare und verlässliche Regulierung von KI im politischen Kontext erforderlich ist, um einen verantwortungsvollen Umgang mit dieser Technologie sicherzustellen. Ein verbindliches Regelwerk sollte dabei sowohl die Chancen berücksichtigen, die KI für politische Kampagnen bietet, als auch dafür Sorge tragen, dass demokratische Prozesse sichergestellt bleiben und grundlegende Freiheitsrechte gewahrt werden (Jungherr, 2024).

Gerade bei KI-generierten Bildern oder Videos ist eine sorgfältige ethische und rechtliche Abwägung erforderlich, um ein Gleichgewicht zwischen Sicherheit und Freiheit zu wahren (Jungherr, 2024; McKay et al., 2024; Taddeo & Floridi, 2018). So braucht es rechtliche Instrumente, mit denen Missbrauch – etwa das Erstellen von Inhalten, die Urheber- oder Persönlichkeitsrechte verletzen, oder das Verbreiten von Desinformationen – wirksam geahndet werden kann. Entsprechende Maßnahmen sollten nicht nur die Verursacher*innen solcher Inhalte betreffen, sondern auch die Anbieter*innen, deren Modelle unzureichende Schutzmechanismen aufweisen (European Commission, 2024). Gleichzeitig müssen Regulierungsansätze sicherstellen, dass die Meinungs- und Kunstfreiheit, wie sie in Artikel 5 des deutschen Grundgesetzes verankert ist, nicht unverhältnismäßig eingeschränkt wird (vgl. Laude & Daum, 2024). Mit Blick auf diese Grundrechte sollte eine Regulierung stets die Frage im Blick behal-

ten, welche Formen künstlerischen Ausdrucks und legitimer Meinungsäußerung weiterhin frei möglich sein müssen.

Kennzeichnung von KI-Inhalten

Ein zentraler Punkt sollte eine klare Verpflichtung sein – insbesondere für Social-Media-Plattformen – KI-generierte Inhalte deutlich erkennbar zu kennzeichnen, damit die Öffentlichkeit deren Ursprung nachvollziehen kann. Hierzu offenbarte unsere Befragung eine breite Unterstützung, die auch mit den regulatorischen Forderungen nach Kennzeichnungspflichten übereinstimmt, wie sie beispielsweise im europäischen AI Act vorgesehen sind (Europäisches Parlament, 2024) und in den Selbstverpflichtungen einiger Parteien und Plattformen festgeschrieben wurden (siehe Kapitel 2.4).

Mit Blick auf die Ergebnisse unseres Experiments muss die Anwendung und Ausführung von Transparenzmaßnahmen wohl durchdacht sein, um eine effektive Aufklärung zu gewährleisten. So erhöht zwar der explizite KI-Hinweis im Bild die Erkennung des KI-Ursprungs, jedoch gilt dies nicht für unseren Aufklärungstext. Daher müssen Kennzeichnungen so gestaltet sein, dass sie deutlich sichtbar sind – andernfalls könnten sie übersehen werden und somit das Bewusstsein für die manipulative Natur solcher Inhalte nicht schärfen (siehe Kapitel 4.3; Friestad & Wright, 1994; Jost et al., 2023). Eine entsprechende Regelung zur Kennzeichnung muss sich damit auseinandersetzen, mit welchen Angaben und ab welchem Grad der KI-Generierung eines Textes, Bildes oder Videos diese gekennzeichnet wer-

den muss (siehe kritische Auseinandersetzung Leibowicz, 2024). Denn macht eine Korrektur eines Textes oder das Retuschieren eines Fotos mittels KI-Tools schon eine Kennzeichnung nötig?

Erweiterung und Durchsetzung von geltendem Recht

Aktuelle Gesetze wie das allgemeine Persönlichkeitsrecht, das Kunsturhebergesetz (KUG) und strafrechtliche Vorschriften bieten zwar gewisse Schutzmechanismen gegen den Missbrauch von KI-generierten Inhalten, stoßen jedoch an ihre Grenzen, da sie nicht speziell auf digitale Manipulationen ausgerichtet sind (Laude & Daum, 2024). Die Identifizierung von Täter*innen gestaltet sich in der anonymen, global vernetzten Online-Welt oft als schwierig. Auch die Datenschutz-Grundverordnung (DSGVO) sowie das Urheberrecht bieten nur begrenzten Schutz gegen die Verbreitung manipulativer Inhalte (Bundesrat, 2024). Politische Werbemaßnahmen unterliegen in Deutschland dem Parteiengesetz und dem Bundeswahlgesetz, die Transparenz und Rechenschaftspflicht von Wahlwerbung vorschreiben. Diese Regelungen greifen jedoch nur bedingt bei digitalen Kampagnen und KI-generierten Inhalten, die oft ohne klare Urheberschaft oder Kennzeichnung verbreitet werden. Dies erleichtert die Manipulation der öffentlichen Meinung und gefährdet die Transparenz des Wahlprozesses (Dobber et al., 2019; G'sell, 2024).

Um diesen Herausforderungen zu begegnen, sollten gezielte Erweiterungen des Rechtsrah-

mens eingeführt werden, darunter spezifische Straftatbestände für die unbefugte Erstellung und Verbreitung von KI-generierten Inhalten. So zielt ein vom Freistaat Bayern im Jahr 2024 in den Bundesrat eingebrachter Gesetzesentwurf darauf ab, einen spezifischen Straftatbestand im Strafgesetzbuch zu schaffen, um den Missbrauch von KI-generierten Medieninhalten zu bekämpfen, die den Anschein authentischer Aufnahmen – insbesondere von Menschen – erwecken (Bundesrat, 2024). Zudem sollten die bestehenden Wahlgesetze so angepasst werden, dass der Einsatz von KI-Technologien in der politischen Werbung ausdrücklich geregelt wird, etwa durch die Einführung einer Kennzeichnungspflicht für KI-generierte Inhalte in politischen Kampagnen. Eine Offenlegungspflicht für den Einsatz von KI-Technologien sowie verstärkte Sanktionen gegen den Missbrauch manipulativer KI-Inhalte könnten zusätzlich zur Transparenz und zur Abschreckung von Missbrauch beitragen.

Auf europäischer Ebene spielt die europäische KI-Verordnung, auch bekannt als AI Act (Europäisches Parlament, 2024), eine zentrale Rolle in der Regulierung von KI-generierten Inhalten und bietet einen strukturierten Rahmen, um die Risiken dieser Technologien besser zu managen. Der AI Act zielt darauf ab, einheitliche Standards für die Entwicklung, den Einsatz und die Überwachung von KI-Systemen innerhalb der EU zu schaffen, wobei die Einstufung der KI-Anwendungen nach Risikokategorien im Mittelpunkt steht. KI-generierte manipulative Inhalte sollten dabei als „hohes Risiko“ eingestuft werden, insbesondere wenn sie das Potenzial haben,

die öffentliche Sicherheit, demokratische Prozesse oder die Meinungsfreiheit zu gefährden. Dies würde strenge Anforderungen an Transparenz, Risikoanalyse und Kennzeichnung solcher Inhalte mit sich bringen (Cupać & Sienknecht, 2024).

Transparenten und verantwortungsbewussten Einsatz von KI auf Social Media durchsetzen

Die Regulierung von Social-Media-Plattformen ist zentral für die Kontrolle der Verbreitung von KI-generierten Inhalten, insbesondere in politischen Kampagnen. Plattformen wie Facebook, X und YouTube haben eine enorme Reichweite und können Einfluss auf die öffentliche Meinung haben. Bestehende Gesetze wie das Netzwerkdurchsetzungsgesetz (NetzDG) in Deutschland und der Digital Services Act (DSA; European Commission, 2021) auf europäischer Ebene reichen jedoch oft nicht aus, um die spezifischen Herausforderungen durch Deepfakes und andere KI-generierte Inhalte zu bewältigen (Sullivan, 2023).

Für eine effektivere Regulierung könnten Social-Media-Plattformen verpflichtet werden, KI-generierte Inhalte mithilfe fortschrittlicher Erkennungstechnologien zu identifizieren und zu kennzeichnen sowie Nutzer*innen über potenziell manipulative Inhalte zu informieren. Darüber hinaus könnten die Algorithmen, die die Verbreitung von Inhalten steuern, angepasst werden, um die Reichweite von Desinformation und Deepfakes zu begrenzen (Bateman & Jackson, 2024). Im Gegensatz zur jüngsten Entscheidung von Social-Media-Plattformen wie Facebook,

Instagram oder X künftig den Wahrheitsgehalt von Inhalten nicht mehr durch unabhängige Organisationen oder Nachrichtenagenturen überprüfen zu lassen (Kaplan, 2025), schlagen wir eine engere Zusammenarbeit mit diesen vor, um die Erkennung und Entfernung manipulativer KI-Inhalte zu verbessern (Klinger & Ohme, 2023). Schließlich sollten Plattformen regelmäßige Audits und Berichte über die Wirksamkeit ihrer Maßnahmen veröffentlichen, um für Transparenz zu sorgen und ihre Glaubwürdigkeit zu erhöhen (Europäisches Parlament, 2024).

Strenge ethische Richtlinien und Transparenzmaßnahmen bei der Nutzung von KI-Tools

Die Regulierung von KI-Plattformen wie OpenAI ist entscheidend, um einen verantwortungsvollen Einsatz von KI-Technologien zu gewährleisten und Missbrauch zu verhindern. Um die potenziellen Risiken von KI-generierten Inhalten zu minimieren, sollten spezifische Regulierungsmaßnahmen eingeführt werden, die Transparenz- und Offenlegungspflichten für KI-Plattformen beinhalten. Dazu gehört die Verpflichtung, Informationen über die Funktionsweise ihrer

Modelle, einschließlich der verwendeten Trainingsdaten und Algorithmen, offenzulegen, um die Auswirkungen und Risiken besser zu verstehen (Europäisches Parlament, 2024; Klinger & Ohme, 2023). Insbesondere bei der Nutzung in politischen Kampagnen sollten umfassende Informationen offengelegt werden, um die Einhaltung der Wahlgesetze zu gewährleisten.

Zudem wäre eine gesetzliche Kennzeichnungspflicht für KI-generierte Inhalte sinnvoll, um Transparenz zu schaffen und den Missbrauch der Technologie zu erschweren beziehungsweise einzudämmen (Jost et al., 2023; Lu & Yuan, 2024). KI-Plattformen sollten auch Erkennungswerkzeuge für KI-generierte Inhalte entwickeln und diese der Öffentlichkeit zugänglich machen. Darüber hinaus sollten Maßnahmen ergriffen werden, die den Einsatz der Modelle zur Erstellung von Desinformation oder manipulativen Inhalten einschränken. Regelmäßige Audits durch unabhängige Stellen könnten die Einhaltung gesetzlicher Vorgaben und ethischer Standards überprüfen und die Rechenschaftspflicht der Plattformen erhöhen (Europäisches Parlament, 2024).

Literatur

Die Links wurden am 13. Januar 2025 zuletzt überprüft.

- AfDnee (2024). Über die Initiative. *AfDnee*. <https://afdnee.de/initiative/>
- Alaga, J., Schuett, J., & Anderljung, M. (2024). *A Grading Rubric for AI Safety Frameworks* (arXiv:2409.08751). arXiv. <https://doi.org/10.48550/arXiv.2409.08751>
- Altay, S., & Gilardi, F. (2024). People are skeptical of headlines labeled as AI-generated, even if true or human-made, because they assume full AI automation. *PNAS nexus*, 3(10). <https://doi.org/10.1093/pnasnexus/pgae403>
- Appel, M., & Prietzel, F. (2022). The detection of political deepfakes. *Journal of Computer-Mediated Communication*, 27(4), zmac008. <https://doi.org/10.1093/jcmc/zmac008>
- Aregger, A. (2023, 4. Juli). *FDP wirbt mit KI-generiertem KlimakleberBild*. Tages-Anzeiger Schweiz. <https://www.tagesanzeiger.ch/fdp-wirbt-mit-ki-generiertem-klimakleber-bild-245669524070>
- Auster, O. (2023, 30. August). *FDP fälscht Wüst-Foto mit Künstlicher Intelligenz*. Kölner Stadt-Anzeiger. <https://www.ksta.de/politik/nrw-politik/streit-im-landtag-fdp-faelscht-wuest-foto-mit-kuenstlicher-intelligenz-636878>
- Auster, O. (2024, 19. März). *SPD schickt für Schulpolitik KI-Kopie von Greta Thunberg ins Rennen*. Kölner Stadt-Anzeiger. <https://www.ksta.de/politik/nrw-politik/schon-wieder-ein-foto-fake-spd-schickt-fuer-schulpolitik-ki-kopie-von-greta-thunberg-ins-rennen-761094>
- Bai, H., Voelkel, J. G., Eichstaedt, J. C., & Willer, R. (2023). *Artificial Intelligence Can Persuade Humans on Political Issues*. OSF preprints. <https://doi.org/10.31219/osf.io/stakv>
- Barari, S., Lucas, C., & Munger, K. (2021). *Political Deepfakes Are As Credible As Other Fake Media And (Sometimes) Real Media*. OSF preprints. <https://doi.org/10.31219/osf.io/cdfh3>
- Bateman, J., & Jackson, D. (2024). *Countering Disinformation Effectively. An Evidence-Based Policy Guide*. Carnegie Endowment for International Peace. <https://carnegieendowment.org/research/2024/01/countering-disinformation-effectively-an-evidence-based-policy-guide>
- Battista, D. (2024). Political communication in the age of artificial intelligence: An overview of deepfakes and their implications. *Society Register*, 8(2), 7–24. <https://doi.org/10.14746/sr.2024.8.2.01>
- Beckert, J., Koch, T., Viererbl, B., & Schulz-Knappe, C. (2021). The disclosure paradox: How persuasion knowledge mediates disclosure effects in sponsored media content. *International Journal of Advertising*, 40(7), 1160–1186. <https://doi.org/10.1080/02650487.2020.1859171>
- Behre, J., Hölig, S., & Möller, J. (2024). Reuters Institute Digital News Report 2024: Ergebnisse für Deutschland. *Arbeitspapiere des Hans-Bredow-Instituts*. <https://doi.org/10.21241/SSOAR.94461>
- Beisch, N., & Koch, W. (2023). ARD/ZDF-Onlinestudie: Weitergehende Normalisierung der Internetnutzung nach Wegfall aller Corona-Schutzmaßnahmen. *Media Perspektiven*, 23, 1–9.

- Bianchi, F., Kalluri, P., Durmus, E., Ladhak, F., Cheng, M., Nozza, D., Hashimoto, T., Jurafsky, D., Zou, J., & Caliskan, A. (2023). Easily Accessible Text-to-Image Generation Amplifies Demographic Stereotypes at Large Scale. *2023 ACM Conference on Fairness, Accountability, and Transparency*, 1493–1504. <https://doi.org/10.1145/3593013.3594095>
- Bieber, C., Heesen, J., Grunwald, A., & Rostalski, F. (2024). *KI im Superwahljahr – Generative KI im Umfeld demokratischer Prozesse* [Whitepaper]. Lernende Systeme. https://doi.org/10.48669/pls_2024-5
- Bimber, B. (2003). *Information and American Democracy: Technology in the Evolution of Political Power*. Cambridge University Press.
- Bitton, D. B., Hoffmann, C. P., & Godulla, A. (2024). Deepfakes in the context of AI inequalities: analysing disparities in knowledge and attitudes. *Information, Communication & Society*, 1–21. <https://doi.org/10.1080/1369118X.2024.2420037>
- Boerman, S. C., van Reijmersdal, E. A., & Neijens, P. C. (2012). Sponsorship Disclosure: Effects of Duration on Persuasion Knowledge and Brand Responses. *Journal of Communication*, 62(6), 1047–1064. <https://doi.org/10.1111/j.1460-2466.2012.01677.x>
- Boerman, S. C., Willemsen, L. M., & Van Der Aa, E. P. (2017). „This Post Is Sponsored“: Effects of Sponsorship Disclosure on Persuasion Knowledge and Electronic Word of Mouth in the Context of Facebook. *Journal of Interactive Marketing*, 38, 82–92. <https://doi.org/10.1016/j.intmar.2016.12.002>
- Bossetta, M., & Schmøkel, R. (2023). Cross-Platform Emotions and Audience Engagement in Social Media Political Campaigning: Comparing Candidates’ Facebook and Instagram Images in the 2020 US Election. *Political Communication*, 40(1), 48–68. <https://doi.org/10.1080/10584609.2022.2128949>
- Brehm, S. S., & Brehm, J. W. (1981). *Psychological reactance: A theory of freedom and control*. Academic Press.
- Breschendorf, F. (2024, 11. August). *KI-Bodybuilder Olaf Scholz auf TikTok: Video könnten ihm „schaden“*. Merkur.de. <https://www.merkur.de/deutschland/ki-bodybuilder-olaf-scholz-auf-tiktok-video-koennten-ihm-schaden-zr-93232561.html>
- Budak, C., Nyhan, B., Rothschild, D. M., Thorson, E., & Watts, D. J. (2024). Misunderstanding the harms of online misinformation. *Nature*, 630(8015), 45–53. <https://doi.org/10.1038/s41586-024-07417-w>
- Bühler, J. (2023). *ChatGPT & Co: Sicherheit von generativer Künstlicher Intelligenz*. TÜV Verband. https://www.tuev-verband.de/fileadmin/user_upload/Content_local/Studien_local/TUEV-Verband_PK_Praesentation_ChatGPT_11_05_2023_final.pdf
- Bundesrat. (2024, 5. Juli). Entwurf eines Gesetzes zum strafrechtlichen Schutz von Persönlichkeitsrechten vor Deepfakes. *Deutscher Bundestag*. <https://dserver.bundestag.de/brd/2024/0222-24B.pdf>
- Burkhardt, S., & Rieder, B. (2024). Foundation models are platform models: Prompting and the political economy of AI. *Big Data & Society*, 11(2), 20539517241247839. <https://doi.org/10.1177/20539517241247839>
- Cao, X., & Kosinski, M. (2024). Large language models know how the personality of public figures is perceived by the general public. *Scientific Reports*, 14(1), Article 1. <https://doi.org/10.1038/s41598-024-57271-z>
- Carrasquillo, A. (2024, 21. August). Democrats use AI in effort to stay ahead with Latino and Black voters. *The Guardian*. <https://www.theguardian.com/us-news/article/2024/aug/21/democrats-ai-black-latino-voters>
- CDU/CSU-Fraktion im Deutschen Bundestag. (2023). *Künstliche Intelligenz als Schlüsseltechnologie für Deutschlands Zukunft* [Positionspapier]. https://www.cducsu.de/sites/default/files/2023-09/PP_K%C3%BCnstliche%20Intelligenz%20als%20Schl%C3%BCsseltechnologie.pdf

- Check Point. (2024, 2. April). Beyond Imagining – How AI is actively used in election campaigns around the world. *Check Point*. <https://blog.checkpoint.com/research/beyond-imagining-how-ai-is-actively-used-in-election-campaigns-around-the-world/>
- Chowdhury, R. (2024, 9. April). *AI-fuelled election campaigns are here – Where are the rules?* nature. <https://www.nature.com/articles/d41586-024-00995-9>
- Clore, G. L., Schwarz, N., & Conway, M. (1994). Affective causes and consequences of social information processing. In R.S.J. Wyer & T.K. Srull (Hrsg.), *Handbook of social cognition* (S. 323–418). Lawrence Erlbaum Associates.
- Cowley, E., & Barron, C. (2008). When product placement goes wrong: The effects of program liking and placement prominence. *Journal of Advertising*, 37(1), 89–98. <https://doi.org/10.2753/JOA0091-3367370107>
- Cupać, J., & Sienknecht, M. (2024). Regulate against the machine: How the EU mitigates AI harm to democracy. *Democratization*, 31(5), 1067–1090. <https://doi.org/10.1080/13510347.2024.2353706>
- Das Zentrum für Politische Schönheit. (2024). Zentrum für Politische Schönheit. <https://politicalbeauty.de/ueber-das-ZPS.html>
- Data Protection Authority of Belgium. (2024). *Artificial Intelligence Systems and the GDPR A Data Protection Perspective*. Data Protection Authority of Belgium. <https://www.gegevensbeschermingsautoriteit.be/publications/information-brochure-on-artificial-intelligence-systems-and-the-gdpr.pdf>
- Der Standard. (2023, 7. September). *Google verlangt Kennzeichnung von KI-Inhalten in politischer Werbung*. Der Standard. <https://www.derstandard.de/story/3000000185996/google-verlangt-kennzeichnung-von-ki-inhalten-in-politischer-werbung>
- DIE LINKE. Sachsen Landesvorstand. (2023). *Umgang mit KI-basierten Tools in der politischen Arbeit und den Wahlkämpfen 2024* (B 8 – 158 – 1; Beschluss des Landesvorstandes vom 30. Juni 2023). https://www.dielinke-sachsen.de/wp-content/uploads/2023/01/B_8_158-1_Umgang_KI.pdf
- DIE LINKE. Thüringen. (2023). *lptAntrag A1: Grundsätze für den Einsatz von KI-Tools in unserer politischen Arbeit*. https://www.die-linke-thueringen.de/fileadmin/LV_Thueringen/dokumente/parteitage/lpt9_tagung1/Beschluesse/lpt9_1_A1.pdf
- Dobber, T., Fathaigh, R. Ó., & Borgesius, F. J. Z. (2019). The regulation of online political micro-targeting in Europe. *Internet Policy Review*, 8(4), 1–20. <https://policyreview.info/articles/analysis/regulation-online-political-micro-targeting-europe>
- Dobber, T., Metoui, N., Trilling, D., Helberger, N., & de Vreese, C. (2021). Do (Microtargeted) Deepfakes Have Real Effects on Political Attitudes? *The International Journal of Press/Politics*, 26(1), 69–91. <https://doi.org/10.1177/1940161220944364>
- Dommett, K. (2023). The 2024 Election Will Be Fought on the Ground, Not By AI. *Political Insight*, 14(4), 4–6. <https://doi.org/10.1177/20419058231218316a>
- Dommett, K., Kefford, G., & Kruschinski, S. (2023). *Data-Driven Campaigning in Political Parties: Five Advanced Democracies Compared*. Oxford University Press.
- Duffy, A., Tandoc, E., & Ling, R. (2020). Too good to be true, too good not to share: The social utility of fake news. *Information, Communication & Society*, 23(13), 1965–1979. <https://doi.org/10.1080/1369118X.2019.1623904>
- Eberl, J.-M., Tolochko, P., Jost, P., Heidenreich, T., & Boomgaarden, H. G. (2020). What’s in a post? How sentiment and issue salience affect users’ emotional reactions on Facebook. *Journal of Information Technology & Politics*, 17(1), 48–65. <https://doi.org/10.1080/19331681.2019.1710318>

- Eiserbeck, A., Maier, M., Baum, J., & Abdel Rahman, R. (2023). Deepfake smiles matter less – The psychological and neural impact of presumed AI-generated faces. *Scientific Reports*, 13(1), 16111. <https://doi.org/10.1038/s41598-023-42802-x>
- Elliott, V. (2024). The WIRED AI Elections Project. *Wired*. <https://www.wired.com/story/generative-ai-global-elections/>
- Europäisches Parlament. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (Artificial Intelligence Act) Text with EEA relevance*. <http://data.europa.eu/eli/reg/2024/1689/oj/eng>
- European Commission. (2021). *Proposal for a Regulation of the European Parliament and of the Council on a Single Market For Digital Services (Digital Services Act) and amending Directive 2000/31/EC*. https://ec.europa.eu/newsroom/dae/document.cfm?doc_id=72148
- European Commission. (2024). *Code of Conduct for the 2024 European Parliament elections | European Commission*. https://commission.europa.eu/document/bebd9b72-fbb9-42f3-bcea-dbace7e0650f_en
- FDP-Bundestagsfraktion. (2023). Selbstverpflichtung der Bundestagsfraktion der freien Demokraten zur verantwortungsvollen Nutzung von KI in der politischen Arbeit. *fdpbt.de*. <https://www.fdpbt.de/sites/default/files/2023-11/selbstverpflichtung-der-bundestagsfraktion-der-freien-demokraten-zur-verantwortungsvollen-nutzung-von-ki-in-ihrer-politischen-arbeit.pdf>
- Feng, S., Park, C. Y., Liu, Y., & Tsvetkov, Y. (2023). *From Pretraining Data to Language Models to Downstream Tasks: Tracking the Trails of Political Biases Leading to Unfair NLP Models* (arXiv:2305.08283). arXiv. <https://doi.org/10.48550/arXiv.2305.08283>
- Festinger, L. (1957). *A theory of cognitive dissonance*. Stanford University Press.
- Fletcher, R., & Nielsen, R. K. (2024). *What does the public in six countries think of generative AI in news? Reuters Institute for the Study of Journalism*. <https://doi.org/10.60625/RISJ-4ZB8-CG87>
- Frank, J., Herbert, F., Ricker, J., Schönherr, L., Eisenhofer, T., Fischer, A., Dürmuth, M., & Holz, T. (2023). *A Representative Study on Human Detection of Artificially Generated Media Across Countries* (arXiv:2312.05976). arXiv. <https://doi.org/10.48550/arXiv.2312.05976>
- Friestad, M., & Wright, P. (1994). The Persuasion Knowledge Model: How People Cope with Persuasion Attempts. *Journal of Consumer Research*, 21(1), 1–31. <https://doi.org/10.1086/209380>
- Friestad, M., & Wright, P. (1995). Persuasion knowledge: Lay people's and researchers' beliefs about the psychology of advertising. *Journal of Consumer Research*, 22(1), 62–74. <https://doi.org/10.1086/209435>
- Früh, H. (2011). *Emotionalisierung durch Nachrichten*. Nomos Verlagsgesellschaft GmbH & Co. KG.
- Fuchs, M. (2024, 2. Januar). *Entwirf mir ein Wahlplakat!* Neue Gesellschaft Frankfurter Hefte. <https://www.frankfurter-hefte.de/artikel/entwirf-mir-ein-wahlplakat-3874/>
- Forgas, J. P. (1995). Mood and judgment: The affect infusion model (AIM). *Psychological Bulletin*, 117(1), 39–66. <https://doi.org/10.1037/0033-2909.117.1.39>
- Geise, S. (2011). *Vision that matters*. VS Verlag für Sozialwissenschaften. <https://doi.org/10.1007/978-3-531-92736-7>
- Geise S., & Lobinger, K. (2015). *Visual Framing. Perspektiven und Herausforderungen der Visuellen Kommunikationsforschung*. Herbert von Halem Verlag.

- Geise, S., Maubach, K., & Boettcher Eli, A. (2024). Picture me in person: Personalization and emotionalization as political campaign strategies on social media in the German federal election period 2021. *New Media & Society*, 14614448231224031. <https://doi.org/10.1177/14614448231224031>
- Geva, M., Goldberg, Y., & Berant, J. (2019). Are We Modeling the Task or the Annotator? An Investigation of Annotator Bias in Natural Language Understanding Datasets. In K. Inui, J. Jiang, V. Ng, & X. Wan (Hrsg.), *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)* (S. 1161–1166). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D19-1107>
- Gillespie, T. (2024). Generative AI and the politics of visibility. *Big Data & Society*, 11(2), 20539517241252131. <https://doi.org/10.1177/20539517241252131>
- Goldstein, J. A., Sastry, G., Musser, M., DiResta, R., Gentzel, M., & Sedova, K. (2023). *Generative Language Models and Automated Influence Operations: Emerging Threats and Potential Mitigations* (arXiv:2301.04246). arXiv. <https://doi.org/10.48550/arXiv.2301.04246>
- Gordon, T. (2024, 11. Januar). *The dawn of the AI election*. Prospect. <https://www.prospectmagazine.co.uk/politics/64396/the-dawn-of-the-ai-election>
- Görmann, M. (2024, 30. August). *Thüringen-Wahl: SPD zieht nach Ärger Video zurück – Machtwort von oben*. Thueringen24. <https://www.thueringen24.de/thueringen/article300387810/thueringen-wahl-maier-spd-fuer.html>
- Groh, M., Sankaranarayanan, A., Singh, N., Kim, D. Y., Lippman, A., & Picard, R. (2024). Human detection of political speech deepfakes across transcripts, audio, and video. *Nature Communications*, 15(1), 7629. <https://doi.org/10.1038/s41467-024-51998-z>
- Gruber, J. B., & Votta, F. (2025). Large Language Models. In *Elgar Encyclopedia of Political Communication*. Edward Elgar Publishing. <https://johannesbgruber.eu/publication/2024-09-06-large-language-models/>
- Grüning, D. J., & Schubert, T. W. (2022). Emotional Campaigning in Politics: Being Moved and Anger in Political Ads Motivate to Support Candidate and Party. *Frontiers in Psychology*, 12, 781851. <https://doi.org/10.3389/fpsyg.2021.781851>
- G'sell, F. (2024). *Regulating under Uncertainty: Governance Options for Generative AI* (SSRN Scholarly Paper 4918704). <https://papers.ssrn.com/abstract=4918704>
- Hackenburg, K., & Margetts, H. (2024). *Evaluating the persuasive influence of political microtargeting with large language models*. *PNAS*, 121(24), e2403116121. <https://doi.org/10.31219/osf.io/wnt8b>
- Hameleers, M., & Van Der Meer, T. G. L. A. (2020). Misinformation and Polarization in a High-Choice Media Environment: How Effective Are Political Fact-Checkers? *Communication Research*, 47(2), 227–250. <https://doi.org/10.1177/0093650218819671>
- Hameleers, M., van der Meer, T. G. L. A., & Dobber, T. (2024). Distorting the truth versus blatant lies: The effects of different degrees of deception in domestic and foreign political deepfakes. *Computers in Human Behavior*, 152, 1–13. <https://doi.org/10.1016/j.chb.2023.108096>
- Harper, A., Gehlen, B., & Ivan, P. (2023, 8. November). *AI use in political campaigns raising red flags into 2024 election*. abc news. <https://abcnews.go.com/Politics/ai-political-campaigns-raising-red-flags-2024-election/story?id=102480464>

- Hase, V., Ausloos, J., Boeschoten, L., Pfiffner, N., Janssen, H., Araujo, T., Carrière, T., Vreese, C. de, Haßler, J., Loecherbach, F., Kmetty, Z., Möller, J., Ohme, J., Schmidbauer, E., Struminskaya, B., Trilling, D., Welbers, K., & Haim, M. (2024). *Fulfilling data access obligations: How could (and should) platforms facilitate data donation studies?* Internet Policy Review: Journal On Internet Regulation, 13(3). <https://policyreview.info/articles/analysis/fulfilling-data-access-obligations>
- Hegewisch, N. (2024, 13. Februar). *König ohne Land*. IPG. <https://www.ipg-journal.de/regionen/asien/artikel/koenig-ohne-land-7310/>
- Huesmann, F. (2024, 1. April). *Wie die AfD mit KI-Bildern Stimmung macht – und wie die anderen Parteien KI nutzen*. RedaktionsNetzwerk Deutschland. <https://www.rnd.de/politik/afd-wie-die-partei-kuenstliche-intelligenz-nutzt-WJIGX6BFQVETPBYPW47YRWXHOM.html>
- Hügelmann, B. (2024, 19. März). Political Prompt Engineering. *Konrad Adenauer Stiftung - Politsnack*. <https://www.kas.de/de/web/politische-bildung/politsnack/detail/-/content/political-prompt-engineering>
- Hugger, K.-U. (2022). Medienkompetenz. In U. Sander, F. von Gross & K.-U. Hugger (Hrsg.). *Handbuch Medienpädagogik* (2. Aufl.) (S. 67–80). Springer VS. https://doi.org/10.1007/978-3-658-23578-9_9
- Huh, J., Nelson, M. R., & Russell, C. A. (2023). ChatGPT, AI Advertising, and Advertising Research and Education. *Journal of Advertising*, 52(4), 447–482. <https://doi.org/10.1080/00913367.2023.2227013>
- Hwang, Y., Ryu, J. Y., & Jeong, S.-H. (2021). Effects of Disinformation Using Deepfake: The Protective Effect of Media Literacy Education. *Cyberpsychology, Behavior, and Social Networking*, 24(3), 188–193. <https://doi.org/10.1089/cyber.2020.0174>
- Iske, S., & Barberi, A. (2022). Medienkompetenz – ein Beipackzettel. *MedienPädagogik: Zeitschrift für Theorie und Praxis der Medienbildung*, 50, 21–46. <https://doi.org/10.21240/mpaed/50/2022.12.02.X>
- Iyer, A., Webster, J., Hornsey, M. J., & Vanman, E. J. (2014). Understanding the power of the picture: The effect of image content on emotional and political responses to terrorism. *Journal of Applied Social Psychology*, 44(7), 511–521. <https://doi.org/10.1111/jasp.12243>
- Jackson, D., & Martin, Z. (2024, 18. Juni). *Forget Deepfakes: Social Listening Might be the Most Consequential Use of Generative AI in Politics*. Tech Policy Press. <https://techpolicy.press/forget-deepfakes-social-listening-might-be-the-most-consequential-use-of-generative-ai-in-politics>
- Jakesch, M., French, M., Ma, X., Hancock, J. T., & Naaman, M. (2019). AI-Mediated Communication: How the Perception that Profile Text was Written by AI Affects Trustworthiness. *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, 1–13. <https://doi.org/10.1145/3290605.3300469>
- Jost, P., Kruschinski, S., Sülflow, M., Haßler, J., & Maurer, M. (2023). Invisible Transparency. How Different Types of Disclaimers on Facebook Affect Whether and How Digital Political Advertising is Perceived. *Policy & Internet* 15(2), 204–222. <https://doi.org/10.1002/poi3.333>
- Jungherr, A. (2016). Four Functions of Digital Tools in Election Campaigns: The German Case. *The International Journal of Press/Politics*, 21(3), 358–377. <https://doi.org/10.1177/1940161216642597>
- Jungherr, A. (2023). Artificial Intelligence and Democracy: A Conceptual Framework. *Social Media + Society*, 9(3), Article 3. <https://doi.org/10.1177/20563051231186353>
- Jungherr, A. (2024). Foundational questions for the regulation of digital disinformation. *Journal of Media Law*, 0(0), 1–10. <https://doi.org/10.1080/17577632.2024.2362484>

- Jungherr, A., & Rauchfleisch, A. (2024). Negative Downstream Effects of Alarmist Disinformation Discourse: Evidence from the United States. *Political Behavior*, 46, 2123–2143. <https://doi.org/10.1007/s11109-024-09911-3>
- Jungherr, A., Rivero, G., & Gayo-Avello, D. (2020). *Retooling Politics: How Digital Media Are Shaping Democracy*. Cambridge University Press. <https://doi.org/10.1017/9781108297820>
- Jungherr, A., & Schroeder, R. (2023). Artificial intelligence and the public arena. *Communication Theory*, 33(2–3), 164–173. <https://doi.org/10.1093/ct/qtad006>
- Kamali, N., Nakamura, K., Chatzimpampas, A., Hullman, J., & Groh, M. (2024). *How to Distinguish AI-Generated Images from Authentic Photographs* (arXiv:2406.08651). arXiv. <https://arxiv.org/abs/2406.08651v1>
- Kaplan, J. (2025, 7. Januar). *More Speech and Fewer Mistakes*. <https://about.fb.com/news/2025/01/meta-more-speech-fewer-mistakes/>
- Kasra, M., Shen, C., & O'Brien, J. F. (2018). Seeing Is Believing: How People Fail to Identify Fake Images on the Web. *Extended Abstracts of the 2018 CHI Conference on Human Factors in Computing Systems*, 1–6. <https://doi.org/10.1145/3170427.3188604>
- Kepplinger, H. M., & Maurer, M. (2005). *Abschied vom rationalen Wähler: Warum Wahlen im Fernsehen entschieden werden* (1. Aufl.). Alber.
- Kero, S., Flaßhoff, G., Mühl, A., Laukötter, E., & Došenović, P. (2023). *Kunst erzeugt KI erzeugt Kunst* (8; Meinungsmonitor Künstliche Intelligenz). Center for Advanced Internet Studies. <https://www.cais-research.de/wp-content/uploads/Factsheet-8-Kultur.pdf>
- Kieslich, K., Dosenovic, P., Starke, C., Lünich, M., & Marcinkowski, F. (2022). *Meinungsmonitor Künstliche Intelligenz. Künstliche Intelligenz im Journalismus. Wie nimmt die Bevölkerung den Einfluss von Künstlicher Intelligenz auf die journalistische Arbeit wahr?* (Factsheet Nr. 4). Center for Advanced Internet Studies, CAIS. <https://www.cais-research.de/wp-content/uploads/Factsheet-4-Journalismus.pdf>
- Kieslich, K., Starke, C., Došenović, P., Keller, B., & Marcinkowski, F. (2020). *Künstliche Intelligenz und Diskriminierung* (2; Meinungsmonitor Künstliche Intelligenz). Center for Advanced Internet Studies. <https://www.cais-research.de/wp-content/uploads/Factsheet-2-Diskriminierung.pdf>
- Klinger, U., & Ohme, J. (2023). *Was die Wissenschaft im Rahmen des Datenzugangs nach Art. 40 DSA braucht: 20 Punkte zu Infrastrukturen, Beteiligung, Transparenz und Finanzierung*. <https://doi.org/10.34669/WI.WPP/8.1>
- Knochel, A. D. (2023). Midjourney Killed the Photoshop Star: Assembling the Emerging Field of Synthography. *Studies in Art Education*, 64(4), 467–481. <https://doi.org/10.1080/00393541.2023.2255085>
- Köbis, N. C., Doležalová, B., & Soraperra, I. (2021). Fooled twice: People cannot detect deepfakes but think they can. *iScience*, 24(11), 103364. <https://doi.org/10.1016/j.isci.2021.103364>
- Kreps, S., McCain, R. M., & Brundage, M. (2022). All the News That's Fit to Fabricate: AI-Generated Text as a Tool of Media Misinformation. *Journal of experimental political science*, 9(1), 104–117. <https://doi.org/10.1017/XPS.2020.37>
- Kruglanski, A. W. (1990). Lay Epistemic Theory in Social-Cognitive Psychology. *Psychological Inquiry*, 1(3), 181–197. https://doi.org/10.1207/s15327965pli0103_1
- Kruschinski, S., Votta, F., Runde, M., Scherer, J., & Schültken, T. (2025). CampAlign Tracker: Eine Plattform zur Transparenz von politischen KI-Inhalten auf Social Media. <https://www.campaigntracker.de/>

- Kühne, R., Wirth, W., & Müller, S. (2012). Der Einfluss von Stimmungen auf die Nachrichtenrezeption und Meinungsbildung: Eine experimentelle Überprüfung des Affect Infusion Models. *M&K Medien & Kommunikationswissenschaft*, 60(3), 414–431. <https://doi.org/10.5771/1615-634x-2012-3-414>
- Kunda, Z. (1990). The case for motivated reasoning. *Psychological Bulletin*, 108(3), 480–498. <https://doi.org/10.1037/0033-2909.108.3.480>
- Łabuz, M., & Nehring, C. (2024). On the way to deep fake democracy? Deep fakes in election campaigns in 2023. *European Political Science*, 23, 454–473. <https://doi.org/10.1057/s41304-024-00482-9>
- LaChapelle, C., & Tucker, C. (2023, 28. November). Generative AI in Political Advertising. *Brennan Center For Justice*. <https://www.brennancenter.org/our-work/research-reports/generative-ai-political-advertising>
- Laude, L., & Daum, A. (2024, 3. Mai). KI als neues Wahlkampfinstrument. *Verfassungsblog*. <https://doi.org/10.59704/9e96c5f69bbfa94c>
- Lauer, S. (2023, 6. April). *Wie gehen AfD & Co mit künstlicher Intelligenz um?* Belltower. <https://www.belltower.news/ki-rechtsaussen-wie-gehen-afd-co-mit-kuenstlicher-intelligenz-um-148183/>
- Leaver, T., & Srdarov, S. (2023). ChatGPT Isn't Magic: The Hype and Hypocrisy of Generative Artificial Intelligence (AI) Rhetoric. *M/C Journal*, 26(5). <https://doi.org/10.5204/mcj.3004>
- Lee, T. B. (2023, 31. Mai). Large language models, explained with a minimum of math and jargon. *Understanding AI*. <https://www.understandingai.org/p/large-language-models-explained-with>
- Leibowicz, C. (2024, 9. September). Lawmakers Push for AI Labels, But Ensuring Media Accuracy Is No Easy Task. *TechPolicy.Press*. <https://techpolicy.press/lawmakers-push-for-ai-labels-but-ensuring-media-accuracy-is-no-easy-task>
- Leonhard, L., & Bartsch, A. (2020). Affektive Wirkungen politischer Kommunikation. In I. Borucki, K. Kleinen-von Königslöw, S. Marschall & T. Zerback (Hrsg.), *Handbuch Politische Kommunikation* (1. Aufl.) (S. 1–17). Springer VS. https://doi.org/10.1007/978-3-658-26242-6_42-1
- Liu, A. C. C., Law, O. M. K., & Law, I. (2022). *Understanding Artificial Intelligence: Fundamentals and Applications* (1. Aufl.). Wiley. <https://doi.org/10.1002/9781119858393>
- Lovato, J., St-Onge, J., Harp, R., Salazar Lopez, G., Rogers, S. P., Haq, I. U., Hébert-Dufresne, L., & Onaolapo, J. (2024). Diverse misinformation: Impacts of human biases on detection of deepfakes on networks. *Npj Complexity*, 1(1), 1–13. <https://doi.org/10.1038/s44260-024-00006-y>
- Lu, H., & Yuan, S. (2024). „I know it's a deepfake“: The role of AI disclaimers and comprehension in the processing of deepfake parodies. *Journal of Communication*, jqae022. <https://doi.org/10.1093/joc/jqae022>
- Masch, L., & Gabriel, O. W. (2020). How Emotional Displays of Political Leaders Shape Citizen Attitudes: The Case of German Chancellor Angela Merkel. *German Politics*, 29(2), 158–179. <https://doi.org/10.1080/09644008.2019.1657096>
- Matthes, J., Schemer, C., & Wirth, W. (2007). More than meets the eye: Investigating the hidden impact of brand placements in television magazines. *International Journal of Advertising*, 26(4), 477–503. <https://doi.org/10.1080/02650487.2007.11073029>
- Maurer, M. (2014). 30 Attitudinal effects in political communication. In C. Reinemann (Hrsg.), *Political Communication* (1. Aufl.) (S. 591–608). DE GRUYTER. <https://doi.org/10.1515/9783110238174.591>
- Maurer, M. (2016). *Nonverbale politische Kommunikation* (1. Aufl.). Springer VS.

- McKay, S., Tenove, C., Gupta, N., Ibañez, J., Mathews, N., & Tworek, H. (2024, 21. August). *Harmful Hallucinations: Generative AI and Elections*. <https://open.library.ubc.ca/soa/cIRcle/collections/facultyresearchandpublications/52383/items/1.0445035>
- Meisoll, A. (2024, 21. März). *FAQ: Wie die Parteien in BW mit KI umgehen*. SWR.de. <https://www.swr.de/swraktuell/baden-wuerttemberg/ki-fake-parteien-afd-bw-100.html>
- MeMo:KI. (2022). Dashboard des Meinungsmonitor Künstliche Intelligenz. CAIS. <https://www.cais-research.de/forschung/memoki/memoki-bevoelkerungsbefragung/>
- MeMo:KI – Medienanalyse. (o. J.). MeMo:KI – Medienanalyse. CAIS. Abgerufen am 9. September 2024, von <https://www.cais-research.de/forschung/memoki/memoki-medienanalyse/>
- Merica, D. (2024, 2. Juni). *Democrats wanted an agreement on using artificial intelligence. It went nowhere*. AP News. <https://apnews.com/article/artificial-intelligence-ai-dnc-campaign-2024-republicans-7c6b78b6a8ded9ad253be9ef491e0284>
- Meta. (2024, 5. April). KI-Inhalte in Meta-Produkten erkennen. Meta. <https://www.meta.com/de-de/help/artificial-intelligence/how-ai-generated-content-is-identified-and-labeled-on-meta>
- Mollick, E. (2024, 12. August). Change blindness. *One Useful Thing*. <https://www.oneusefulthing.org/p/change-blindness>
- Mühlenkamp, M. (2024 21. November). Welche Rolle spielt KI im Wahlkampf? Abgerufen am 19.12.2024, von <https://www.tagesschau.de/inland/gesellschaft/kuenstliche-intelligenz-wahlkampf-100.html>
- Müllender, M. (2024, 22. April). *Wenn Robo-Scholz anruft*. taz.de. <https://taz.de/Kuenstliche-Intelligenz-im-Wahlkampf/!6005209/>
- Nabi, R. L. (2003). Exploring the Framing Effects of Emotion: Do Discrete Emotions Differentially Influence Information Accessibility, Information Seeking, and Policy Preference? *Communication Research*, 30(2), 224–247. <https://doi.org/10.1177/0093650202250881>
- Nelson, M. R., Ham, C. D., & Haley, E. (2021). What Do We Know about Political Advertising? Not Much! Political Persuasion Knowledge and Advertising Skepticism in the United States. *Journal of Current Issues & Research in Advertising*, 42(4), 329–353. <https://doi.org/10.1080/10641734.2021.1925179>
- Neu, V. (2024). *Die digitale Spaltung der Gesellschaft: Ergebnisse aus einer repräsentativen Umfrage zu Künstlicher Intelligenz* (Monitor Wahl- und Sozialforschung). KAS. <https://www.kas.de/documents/252038/29391852/Die+digitale+Spaltung+der+Gesellschaft.pdf/8570be55-a7ef-849e-3c76-310e10fb2c44?version=1.0&t=1710321957363>
- Nightingale, S. J., & Farid, H. (2022). AI-synthesized faces are indistinguishable from real faces and more trustworthy. *Proceedings of the National Academy of Sciences*, 119(8), e2120481119. <https://doi.org/10.1073/pnas.2120481119>
- Novelli, C., & Sandri, G. (2024). *Digital Democracy in the Age of Artificial Intelligence* (SSRN Scholarly Paper 4901264). <https://papers.ssrn.com/abstract=4901264>
- OMF. (2024). Die 6 besten KI-generierten Werbekampagnen bekannter Marken. OMF. <https://omf.ai/blog/top-ki-werbung-bekannter-marken/>
- OpenAI. (2016, 16. Juni). Generative models. OpenAI. <https://openai.com/index/generative-models/>

- OpenAI. (2024a, 14. Mai). How OpenAI is approaching 2024 worldwide elections. *OpenAI*. <https://openai.com/index/how-openai-is-approaching-2024-worldwide-elections/>
- OpenAI. (2024b, 30. Mai). Disrupting deceptive uses of AI by covert influence operations. *OpenAI*. <https://openai.com/index/disrupting-deceptive-uses-of-AI-by-covert-influence-operations/>
- Oppenlaender, J. (2024). A Taxonomy of Prompt Modifiers for Text-To-Image Generation. *Behaviour & Information Technology*, 43 (15), 3763–3776. <https://doi.org/10.1080/0144929X.2023.2286532>
- Petty, R. E., & Cacioppo, J. T. (1986). The Elaboration Likelihood Model of Persuasion. *Advances in Experimental Social Psychology* 19, 123–205. [https://doi.org/10.1016/S0065-2601\(08\)60214-2](https://doi.org/10.1016/S0065-2601(08)60214-2)
- Phoenix, J., & Taylor, M. (2024). *Prompt Engineering for Generative AI: Future-Proof Inputs for Reliable AI Outputs* (1. Aufl.). O'Reilly Media.
- Prosser, R. (2024, 23. August). *Elon Musk is turned into an armed robber using his own AI system Grok*. Mail Online. <https://www.dailymail.co.uk/news/article-13772095/Elon-Musk-armed-robber-AI-Grok-deepfake-video-Harris-Trump-Obama-Biden-Putin.html>
- Reveland, C., & Siggelkow, P. (2023, 13. November). *Künstliche Intelligenz: Falsche tagesschau-Audiodateien im Umlauf*. tagesschau.de. <https://www.tagesschau.de/faktenfinder/tagesschau-audio-fakes-100.html>
- Robertson, A. (2024, 14. August). *X's new AI image generator will make anything from Taylor Swift in lingerie to Kamala Harris with a gun*. The Verge. <https://www.theverge.com/2024/8/14/24220173/xai-grok-image-generator-misinformation-offensive-images>
- Robins-Early, N. (2024, 19. August). *Trump posts deepfakes of Swift, Harris and Musk in effort to shore up support*. The Guardian. <https://www.theguardian.com/us-news/article/2024/aug/19/trump-ai-swift-harris-musk-deepfake-images>
- Roelcke, T. (2023, 13. November). „Euer Zuhause. Unser Auftrag“: Berliner Senat wirbt mit KI-Gesichtern um Akzeptanz fürs Bauen. Tagesspiegel. <https://www.tagesspiegel.de/berlin/berliner-wirtschaft/euer-zuhause-unser-auftrag-berliner-senat-wirbt-fur-mehr-akzeptanz-furs-bauen-10772335.html>
- Rottach, M. (2024, 21. März). *KI: AfD Göppingen wirbt mit einem Fake-Gesicht*. SWR.de. <https://www.swr.de/swraktuell/baden-wuerttemberg/stuttgart/afd-deepfakes-wahlkampf-100.html>
- Ryan, M. (2020). In AI We Trust: Ethics, Artificial Intelligence, and Reliability. *Science and Engineering Ethics*, 26(5), 2749–2767. <https://doi.org/10.1007/s11948-020-00228-y>
- Sætra, H. S. (2023). Generative AI: Here to stay, but for good? *Technology in Society*, 75, 102372. <https://doi.org/10.1016/j.techsoc.2023.102372>
- Sandri, G., Lupato, F. G., Meloni, M., von Nostitz, F., & Barberà, O. (2024). Mapping the digitalisation of European political parties. *Information, Communication & Society*, 0(0), 1–22. <https://doi.org/10.1080/1369118X.2024.2343369>
- Sanguinetti, P., & Palomo, B. (2024). An Alien in the Newsroom: AI Anxiety in European and American Newspapers. *Social Sciences*, 13(11), 608. <https://doi.org/10.3390/socsci13110608>
- Sanseviero, O., Cuenca, P., Passos, A., & Whitaker, J. (2025). *Hands-On Generative AI with Transformers and Diffusion Models* (1. Aufl.). O'Reilly Media.

- Sastry, G., Heim, L., Belfield, H., Anderljung, M., Brundage, M., Hazell, J., O’Keefe, C., Hadfield, G. K., Ngo, R., Pilz, K., Gor, G., Bluemke, E., Shoker, S., Egan, J., Trager, R. F., Avin, S., Weller, A., Bengio, Y., & Coyle, D. (2024). *Computing Power and the Governance of Artificial Intelligence* (arXiv:2402.08797). arXiv. <https://doi.org/10.48550/arXiv.2402.08797>
- Sato, M. (2025, 3. Januar). Meta’s AI-generated bot profiles are not being received well. <https://www.theverge.com/2025/1/3/24334946/meta-ai-profiles-instagram-facebook-bots>
- Schemer, C. (2010). Der affektive Einfluss von politischer Werbung in Kampagnen auf Einstellungen. *M&K Medien & Kommunikationswissenschaft*, 58(2), 227–246.
- Schindler, S., Zell, E., Botsch, M., & Kissler, J. (2017). Differential effects of face-realism and emotion on event-related brain potentials and their implications for the uncanny valley theory. *Scientific Reports*, 7(1), 45003. <https://doi.org/10.1038/srep45003>
- Schlude, A., Schwind, M., Mendel, U., Stürz, R. A., Harles, D., & Fischer, M. (2023). *Verbreitung und Akzeptanz generativer KI in Deutschland und an deutschen Arbeitsplätzen*. Bidt DE. <https://doi.org/10.35067/xypq-kn69>
- Schmitt-Beck, R. (2000). *Politische Kommunikation und Wählerverhalten* (1. Aufl.). VS Verlag für Sozialwissenschaften. <https://doi.org/10.1007/978-3-322-80381-8>
- Schnöller, J. (2024). *Agile Demokratie: Wie Künstliche Intelligenz bessere Politik ermöglicht* (1. Aufl.). Murmann.
- Scholl, M. (2024, 27. August). *Landtagswahlen 2024: So hilft den Parteien KI im Wahlkampf*. ZDFheute. <https://www.zdf.de/nachrichten/politik/deutschland/ki-kuenstliche-intelligenz-landtagswahl-parteien-100.html>
- Schubert, M. (2023, 3. Dezember). *KI in der Politik: Gefährlich für Demokratie? Experte warnt vor „Brandbeschleuniger“*. Der Westen. <https://www.derwesten.de/politik/kuenstliche-intelligenz-ki-olaf-scholz-fake-desinformation-gruene-id300739601.html>
- Schueler, M., Romano, S., Stanusch, N., Çetin, R. B., Tabiti, S., Faddoul, M., & Lilley, I. (2024). Artificial Elections. Exposing the Use of Generative AI Imagery in the Political Campaigns of the 2024 French Elections. *AI Forensics*. https://aiforensics.org/uploads/Report_Artificial_Elections_81d14977e9.pdf
- Shah, F. (2024, 30. April). *AI companies are making millions producing election content in India*. Rest of world. <https://restofworld.org/2024/india-elections-ai-content/>
- Shen, C., Kasra, M., Pan, W., Bassett, G. A., Malloch, Y., & O’Brien, J. F. (2019). Fake images: The effects of source, intermediary, and digital media literacy on contextual assessment of image credibility online. *New Media & Society*, 21(2), 438–463. <https://doi.org/10.1177/1461444818799526>
- Short, C. E., & Short, J. C. (2023). The artificially intelligent entrepreneur: ChatGPT, prompt engineering, and entrepreneurial rhetoric creation. *Journal of Business Venturing Insights*, 19, e00388. <https://doi.org/10.1016/j.jbvi.2023.e00388>
- Sindermann, C., Cooper, A., & Montag, C. (2020). A short review on susceptibility to falling for fake political news. *Current Opinion in Psychology*, 36, 44–48. <https://doi.org/10.1016/j.copsyc.2020.03.014>
- SPD-Bundestagsfraktion. (2024, 2. Juli). *Leitlinien zur Nutzung von generativer KI in der parlamentarischen Arbeit*. <https://www.spdfraktion.de/system/files/documents/position-ki-leitlinien.pdf>
- Spiegel.de (2024, 22. Dezember). *Parteien beschließen Fairnessabkommen für den Wahlkampf*. <https://www.spiegel.de/politik/deutschland/bundestagswahl-2025-parteien-beschliessen-fairnessabkommen-fuer-den-wahlkampf-a-dbc46911-934c-48be-bd65-7b54e16d1487>

- Sullivan, D. (2023, 19. Juli). *Unpacking “Systemic Risk” Under the EU’s Digital Service Act*. Tech Policy Press. <https://techpolicy.press/unpacking-systemic-risk-under-the-eus-digital-service-act/>
- Sundar, S. S. (2008). The MAIN Model: A Heuristic Approach to Understanding Technology Effects on Credibility. In M. J. Metzger & A. J. Flanagin (Hrsg.), *Digital Media, Youth, and Credibility* (S. 73–100). MIT Press. <https://doi.org/10.1162/dmal.9780262562324.073>
- Taddeo, M., & Floridi, L. (2018). How AI can be a force for good. *Science*, 361(6404), 751-752. <https://doi.org/10.1126/science.aat5991>
- Tahir, R., Batool, B., Jamshed, H., Jameel, M., Anwar, M., Ahmed, F., Zaffar, M. A., & Zaffar, M. F. (2021). Seeing is Believing: Exploring Perceptual Differences in DeepFake Videos. *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, 1–16. <https://doi.org/10.1145/3411764.3445699>
- Tardáguila, C. (2024, 18. März). *New Analysis Reveals Scope of „Fake News“ Referencing or Produced by AI in Brazil; Little Related to Elections or Democracy, For Now | TechPolicy.Press*. Tech Policy Press. <https://techpolicy.press/new-analysis-reveals-scope-of-fake-news-referencing-or-produced-by-ai-in-brazil-little-related-to-elections-or-democracy-for-now>
- t-online. (2024, 14. Juni). *KI-Video: Wüst begrüßt Fußballfans in elf Sprachen*. t-online. https://www.t-online.de/region/duesseldorf/id_100427474/em-hendrik-wuest-gruesst-fussball-fans-in-elf-sprachen-dank-ki.html
- Töpfer, E. (2024, 6. August). *Wirbel um CDU-Plakate: Steckt Künstliche Intelligenz dahinter?* TAG24. <https://www.tag24.de/nachrichten/politik/deutschland/wahlen/landtagswahl-sachsen/kuenstliche-intelligenz-in-cdus-plakatkampagne-entdeckt-3305836>
- Tucciarelli, R., Vehar, N., Chandaria, S., & Tsakiris, M. (2022). On the realness of people who do not exist: The social processing of artificial faces. *iScience*, 25(12), 105441. <https://doi.org/10.1016/j.isci.2022.105441>
- Tunstall, L., Werra, L. von, Wolf, T., & Géron, A. (2023). *Natural Language Processing mit Transformern: Sprachanwendungen mit Hugging Face erstellen*. O'Reilly.
- Universum. (2024). Superwahljahr 2024: KI erobert den Wahlkampf. *Universum*. <https://www.universum.com/superwahljahr-2024-ki-erobert-den-wahlkampf>
- Urman, A., & Makhortykh, M. (2023). *The Silence of the LLMs: Cross-Lingual Analysis of Political Bias and False Information Prevalence in ChatGPT, Google Bard, and Bing Chat*. <https://doi.org/10.31219/osf.io/q9v8f>
- Vaccari, C., & Chadwick, A. (2020). Deepfakes and Disinformation: Exploring the Impact of Synthetic Political Video on Deception, Uncertainty, and Trust in News. *Social Media + Society*, 6(1), 2056305120903408. <https://doi.org/10.1177/2056305120903408>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2023). *Attention Is All You Need* (arXiv:1706.03762). arXiv. <https://doi.org/10.48550/arXiv.1706.03762>
- Vogler, D., Eisenegger, M., Fürst, S., Udris, L., Ryffel, Q., Rivière, M., & Schäfer, M. S. (2023). Künstliche Intelligenz in der journalistischen Nachrichtenproduktion: Wahrnehmung und Akzeptanz in der Schweizer Bevölkerung. In Forschungszentrum Öffentlichkeit und Gesellschaft (Hrsg.), *Jahrbuch Qualität der Medien* (S. 33–46). Schwabe. <https://doi.org/10.5167/UZH-235608>
- Votta, F., & de León, E. (2024). *AlgoSoc AI Opinion Monitor*. AlgoSoc AI Opinion Monitor. <https://monitor.algosoc.org/en/>

- Wang, C., Boerman, S. C., Kroon, A. C., Möller, J., & H de Vreese, C. (2024). The artificial intelligence divide: Who is the most vulnerable? *New Media & Society*, 14614448241232345. <https://doi.org/10.1177/14614448241232345>
- Wankhede, C. (2024, 6. März). What is Midjourney AI and how does it work? *Android Authority*. <https://www.androidauthority.com/what-is-midjourney-3324590/>
- Warrlich, S. (2023, 20. November). *Wer trägt die Verantwortung im Falle des Scheiterns?* Stuttgarter Zeitung. <https://www.stuttgarter-zeitung.de/inhalt.kolumne-wer-wird-abgesaegt.eb69051c-8050-4eef-8efd-affab15ca212.html>
- Weikmann, T., Greber, H., & Nikolaou, A. (2024). After Deception: How Falling for a Deepfake Affects the Way We See, Hear, and Experience Media. *The International Journal of Press/Politics*, 19401612241233539. <https://doi.org/10.1177/19401612241233539>
- Wendt, R., Riesmeyer, C., Leonhard, L., Hagner, J., & Kühn, J. (2024). *Algorithmen und Künstliche Intelligenz im Alltag von Jugendlichen: Forschungsbericht für die Bayerische Landeszentrale für neue Medien*. Bayerische Landeszentrale für neue Medien (BLM). https://www.blm.de/files/pdf2/blm_schriftenreihe_111.pdf
- Wigger, K. (2025, 5. Januar). *OpenAI is losing money on its pricey ChatGPT Pro plan*, CEO Sam Altman says. TechCrunch.com. <https://techcrunch.com/2025/01/05/openai-is-losing-money-on-its-pricey-chatgpt-pro-plan-ceo-sam-altman-says/>
- Williams, K. (2024, 12. September). *US States Struggle to Define „Deepfakes“ and Related Terms as Technically Complex Legislation Proliferates*. Tech Policy Press. <https://techpolicy.press/us-states-struggle-to-define-deepfakes-and-related-terms-as-technically-complex-legislation-proliferates>
- Wolfram, S. (2023). *What is ChatGPT doing ... And why does it work?* Wolfram Media, Inc.
- Wodecki, B. (2023). *DeSantis Uses AI Images to Attack Trump Record on Fauci*. aibusiness.com <https://aibusiness.com/responsible-ai/desantis-uses-ai-images-to-attack-trump-record-on-fauci>
- World Economic Forum. (2024). *The Global Risks Report 2024*. https://www3.weforum.org/docs/WEF_The_Global_Risks_Report_2024.pdf
- Wu, H. D. (2024). Physiological Response to Political Advertisement: Examining the Influence of Partisan and Issue Congruence on Attention and Emotion. *International Journal of Communication*, 18, 1870–1890. <https://ijoc.org/index.php/ijoc/article/view/20686/4544>

Verzeichnis der Abbildungen

Abbildung 1:	Struktur für das Prompt Engineering in <i>Midjourney</i>	15
Abbildung 2:	KI-generiertes Bild von Donald Trump und Dr. Anthony Fauci (links) und KI-generiertes Überwachungsvideo von Kamala Harris beim Begehen einer Straftat (rechts).....	17
Abbildung 3:	KI-generierte Avatare von Prabowo Subinato und seines „Running Mate“ Gibran Rakabuming Raka (links) und KI-generierte Social-Media-Posts der Reconquête (rechts)	18
Abbildung 4:	KI-generiertes Bild der SPD-Landtagsfraktion NRW	19
Abbildung 5:	KI-generierte Social-Media-Posts der CDU Sachsen (links) und der FDP-Landtagsfraktion NRW (rechts)	20
Abbildung 6:	KI-generierte Social-Media-Posts der AfD-Fraktion Baden-Württemberg (oben links), der AfD-Fraktion Rheinland-Pfalz (oben rechts) und Screenshots aus KI-generierten Video der AfD Brandenburg (unten)	21
Abbildung 7:	KI-generierter Social-Media-Post der Linken Sachsen	22
Abbildung 8:	KI-generierte Bilder der Kampagne „AfD-Verbot“ vom Zentrum für politische Schönheit (oben) sowie der Kampagne „AfDnee“ (unten)	23
Abbildung 9:	Transparenzhinweis für die KI-Nutzung auf Instagram in unbezahlten Posts (links) und bezahlter Werbung (rechts)	28
Abbildung 10:	Transparenzhinweis für die KI-Nutzung durch DALL-E	29
Abbildung 11:	Allgemeines Wissen zu KI	41
Abbildung 12:	Aufbau des Fragebogens	43
Abbildung 13:	Kenntnis von KI in der Politik	44
Abbildung 14:	Bewertung von KI in der Politik	45
Abbildung 15:	Einschätzung von Parteien, die KI einsetzen.....	46
Abbildung 16:	Akzeptanz von KI in der Politik nach Einsatzgebieten.....	47
Abbildung 17:	Chancen eines KI-Einsatzes in der Politik	48
Abbildung 18:	Risiken eines KI-Einsatzes in der Politik	48
Abbildung 19:	Regulierung des Einsatzes von KI in der Politik	50
Abbildung 20:	Untersuchungsdesign und -ablauf des Online-Experiments zur Wahrnehmung und Wirkung von KI-generierten Kampagnenbildern	66
Abbildung 21:	Erkennung KI-generierter Bilder abhängig von tatsächlicher KI-Generierung, Valenz, Aufklärungstext und Hinweislablel.....	68
Abbildung 22:	Durch die Kampagnenbilder ausgelöste Emotionen abhängig von tatsächlicher KI-Generierung, Valenz, Aufklärungstext und Hinweislablel im Bild	69
Abbildung 23:	Bewertung der Kampagnenbilder in ihrer Natürlichkeit, Überzeugungskraft und Aufdringlichkeit abhängig von tatsächlicher KI-Generierung, Valenz, Aufklärungstext und Hinweislablel.....	70
Abbildung 24:	Einfluss KI-generierter Kampagnenbilder auf die Bewertung des KI-Einsatzes, der Partei und von KI als Gefahr	72

Hinweise zu den Autor*innen

Dr. Simon Kruschinski ist Forschungsprojektleiter und Postdoktorand am Lehrbereich Politische Kommunikation des Instituts für Publizistik an der Johannes Gutenberg-Universität Mainz. In seiner Forschung geht er schwerpunktmäßig der Frage nach, wie mit digitalen Kommunikationsstrategien und Technologien Menschen oder Demokratien beeinflusst werden können.

Dr. Pablo Jost ist Kommunikationswissenschaftler am Lehrbereich Politische Kommunikation des Instituts für Publizistik der Johannes Gutenberg-Universität Mainz. Aktuell forscht und lehrt er als Gastprofessor am Institut für Journalistik und Kommunikationsforschung an der HMTM Hannover. Zu seinen Forschungsinteressen gehören die mediale Darstellung gesellschaftlicher Kontroversen, die digitale Mobilisierung radikaler Protestgruppen sowie die Kommunikation politischer Akteur*innen und deren Anpassung an die Bedingungen der Digitalisierung.

Tobias Scherer ist wissenschaftlicher Mitarbeiter am Lehrbereich Politische Kommunikation des Instituts für Publizistik an der Johannes Gutenberg-Universität Mainz. Schwerpunktmäßig forscht er im Gebiet der Medien- und Informationskompetenz sowie zu Desinformation.

Hannah Fecher ist wissenschaftliche Mitarbeiterin am Lehrbereich Politische Kommunikation des Instituts für Publizistik an der Johannes Gutenberg-Universität Mainz. Ihre Forschungsinteressen umfassen den Einsatz und die Wirkung neuer Technologien in der politischen Kommunikation.

- Nr. 74 Tragische Einzelfälle? Trendreport zur Berichterstattung über Gewalt gegen Frauen (Christine E. Meltzer)
- Nr. 73 Social-Media-Partei AfD? Digitale Landtagswahlkämpfe im Vergleich (Maik Fielitz, Harald Sick, Michael Schmidt, Christian Donner)
- Nr. 72 Öffentlichkeit ohne Journalismus? Rollenverschiebungen im lokalen Raum (Barbara Witte, Gerhard Syben)
- Nr. 71 Finanzbildung als politisches Projekt. Eine kritische Analyse der FDP-Initiative zur finanziellen Bildung (Thomas Höhne)
- Nr. 70 ‚Falsche Propheten‘ in Sachsen. Extrem rechte Agitation im Landtag (Ulf Bohmann, Moritz Heinrich, Matthias Sommer)
- Nr. 69 ARD, ZDF und DLR im Wandel. Reformideen und Zukunftsperspektiven (Jan Christopher Kalbhenn)
- Nr. 68 Engagiert und gefährdet. Ausmaß und Ursachen rechter Bedrohungen der politischen Bildung in Sachsen (Thomas Laux, Teresa Lindenauer)
- Nr. 67 Viel Kraft – wenig Biss. Wirtschaftsberichterstattung in ARD und ZDF (Henrik Müller, Gerret von Nordheim)
- Nr. 66 Reklame für Klimakiller. Wie Fernseh- und YouTube-Werbung den Medienstaatsvertrag verletzt (Uwe Krüger, Katharina Forstmair, Alexandra Hilpert, Laurie Stührenberg)
- Nr. 65 Schlecht beraten? Die wirtschaftspolitischen Beratungsgremien der Bundesregierung in der Kritik (Dieter Plehwe, Moritz Neujeffski, Jürgen Nordmann)
- Nr. 64 Arbeitswelt und Demokratie in Ostdeutschland. Erlebte Handlungsfähigkeit im Betrieb und (anti)demokratische Einstellungen (Johannes Kiess, Alina Wesser-Saalfrank, Sophie Bose, Andre Schmidt, Elmar Brähler & Oliver Decker)
- Nr. 63 Konzerne im Klimacheck. ‚Integrated Business Reporting‘ als neuer Ansatz der Unternehmensberichterstattung (Lutz Frühbrodt)
- Nr. 62 Auf der Suche nach Halt. Die Nachwendegeneration in Krisenzeiten (Simon Storks, Rainer Faus, Jana Faus)
- Nr. 61 Desiderius-Erasmus-Stiftung. Immer weiter nach rechts außen (Arne Semsrott, Matthias Jakubowski)
- Nr. 60 Vom Winde verdreht? Mediale Narrative über Windkraft, Naturschutz und Energiewandel (Georgiana Banita)
- Nr. 59 Radikalisiert und etabliert. Die AfD vor dem Superwahljahr 2024 (Wolfgang Schroeder, Bernhard Weißels)
- Nr. 58 Antisemitismus. Alte Gefahr mit neuen Gesichtern (Michael Kraske)

Die Otto Brenner Stiftung ...

... ist die gemeinnützige Wissenschaftsstiftung der IG Metall. Sie hat ihren Sitz in Frankfurt am Main. Als Forum für gesellschaftliche Diskurse und Einrichtung der Forschungsförderung ist sie dem Ziel der sozialen Gerechtigkeit verpflichtet. Besonderes Augenmerk gilt dabei dem Ausgleich zwischen Ost und West.

... initiiert den gesellschaftlichen Dialog durch Veranstaltungen, Workshops und Kooperationsveranstaltungen (z. B. im Herbst die OBS-Jahrestagungen), organisiert Konferenzen, lobt jährlich den „Otto Brenner Preis für kritischen Journalismus“ aus, fördert wissenschaftliche Untersuchungen zu sozialen, arbeitsmarkt- und gesellschaftspolitischen Themen und legt aktuelle medienkritische und -politische Analysen vor.

... informiert regelmäßig mit einem Newsletter über Projekte, Publikationen, Termine und Veranstaltungen.

... veröffentlicht die Ergebnisse ihrer Forschungsförderung in der Reihe „OBS-Arbeitshefte“ oder als Arbeitspapiere (nur online). Die Arbeitshefte werden, wie auch alle anderen Publikationen der OBS, kostenlos abgegeben. Über die Homepage der Stiftung können sie auch elektronisch bestellt werden. Vergriffene Hefte halten wir als PDF zum Download bereit unter: www.otto-brenner-stiftung.de/wissenschaftsportalk/publikationen/

... freut sich über jede ideelle Unterstützung ihrer Arbeit. Aber wir sind auch sehr dankbar, wenn die Arbeit der OBS materiell gefördert wird.

... ist zuletzt durch Bescheid des Finanzamtes Frankfurt am Main V (-Höchst) vom 16. November 2023 als ausschließlich und unmittelbar gemeinnützig anerkannt worden. Aufgrund der Gemeinnützigkeit der Otto Brenner Stiftung sind Spenden steuerlich absetzbar bzw. begünstigt.

Unterstützen Sie unsere Arbeit, z. B. durch eine zweckgebundene Spende

Spenden erfolgen nicht in den Vermögensstock der Stiftung, sie werden ausschließlich und zeitnah für die Durchführung der Projekte entsprechend dem Verwendungszweck genutzt.

Bitte nutzen Sie folgende Spendenkonten:

Für Spenden mit zweckgebundenem Verwendungszweck zur Förderung von Wissenschaft und Forschung zum Schwerpunkt:

- Förderung der internationalen Gesinnung und des Völkerverständigungsgedankens

Bank: HELABA Frankfurt/Main
IBAN: DE11 5005 0000 0090 5460 03
BIC: HELA DE FF

Für Spenden mit zweckgebundenem Verwendungszweck zur Förderung von Wissenschaft und Forschung zu den Schwerpunkten:

- Angleichung der Arbeits- und Lebensverhältnisse in Ost- und Westdeutschland (einschließlich des Umweltschutzes)
- Entwicklung demokratischer Arbeitsbeziehungen in Mittel- und Osteuropa
- Verfolgung des Zieles der sozialen Gerechtigkeit

Bank: HELABA Frankfurt/Main
IBAN: DE86 5005 0000 0090 5460 11
BIC: HELA DE FF

Geben Sie bitte Ihre vollständige Adresse auf dem Überweisungsträger an, damit wir Ihnen nach Eingang der Spende eine Spendenbescheinigung zusenden können. Oder bitten Sie in einem kurzen Schreiben an die Stiftung unter Angabe der Zahlungsmodalitäten um eine Spendenbescheinigung. Verwaltungsrat und Geschäftsführung der Otto Brenner Stiftung danken für die finanzielle Unterstützung und versichern, dass die Spenden ausschließlich für den gewünschten Verwendungszweck genutzt werden.

Aktuelle Ergebnisse der Forschungsförderung in der Reihe „OBS-Arbeitshefte“

■ OBS-Arbeitsheft 114*

Marlis Prinzing, Mira Keßler, Melanie Radue

Berichten über Leid und Katastrophen

Die Ahrtalflut 2021 aus Betroffenen- und Mediensicht sowie Lehren für künftige Krisen

■ OBS-Arbeitsheft 113*

Janine Greyer-Stock, Julia Lück-Benz

Moderne Wirtschaftsberichterstattung?

Wie Podcasts auf Spotify und in der ARD Audiothek über Wirtschaft sprechen

■ OBS-Arbeitsheft 112*

Leif Kramp, Stephan Weichert

Whitepaper Non-Profit-Journalismus

Handreichungen für Medien, Politik und Stiftungswesen

■ OBS-Arbeitsheft 111*

Janis Brinkmann

Journalistische Grenzgänger

Wie die Reportage-Formate von funk Wirklichkeit konstruieren

■ OBS-Arbeitsheft 110*

Henning Eichler

Journalismus in sozialen Netzwerken

ARD und ZDF im Bann der Algorithmen?

■ OBS-Arbeitsheft 109*

Barbara Witte, Gerhard Syben

Erosion von Öffentlichkeit

Freie Journalist*innen in der Corona-Pandemie

■ OBS-Arbeitsheft 108*

Victoria Sophie Teschendorf, Kim Otto

Framing in der Wirtschaftsberichterstattung

Der EU-Italien-Streit 2018 und die Verhandlungen über Corona-Hilfen 2020 im Vergleich

■ OBS-Arbeitsheft 107*

Leif Kramp, Stephan Weichert

Konstruktiv durch Krisen?

Fallanalysen zum Corona-Journalismus

■ OBS-Arbeitsheft 106*

Lutz Frühbrodt, Ronja Auerbacher

Den richtigen Ton treffen

Der Podcast-Boom in Deutschland

■ OBS-Arbeitsheft 105*

Hektor Haarkötter, Filiz Kalmuk

Medienjournalismus in Deutschland

Seine Leistungen und blinden Flecken

■ OBS-Arbeitsheft 104*

Valentin Sagvosdkin

Qualifiziert für die Zukunft?

Zur Pluralität der wirtschaftsjournalistischen Ausbildung in Deutschland

* Printfassung leider vergriffen; Download weiterhin möglich.

Diese und weitere Publikationen der OBS finden Sie unter www.otto-brenner-stiftung.de
Otto Brenner Stiftung | Wilhelm-Leuschner-Straße 79 | D-60329 Frankfurt/Main

OBS-Arbeitspapier 75

Künstliche Intelligenz in politischen Kampagnen

Akzeptanz, Wahrnehmung und Wirkung