

Secondary Publication



Jungherr, Andreas; Rauchfleisch, Adrian; Wuttke, Alexander

Artificial Intelligence in Election Campaigns : Perceptions, Penalties, and Implications

Date of secondary publication: 06.07.2026

Version of Record (Published Version), Article

Persistent identifier: urn:nbn:de:bvb:473-irb-115966x

Primary publication

Jungherr, Andreas; Rauchfleisch, Adrian; Wuttke, Alexander (2026): Artificial Intelligence in Election Campaigns : Perceptions, Penalties, and Implications, in: Political Communication, London [u.a.]: Routledge, Taylor & Francis Group, Vol. 43, No. 4, pp. 585–606, doi: 10.1080/10584609.2025.2611913.

Legal Notice

This work is protected by copyright and/or the indication of a licence. You are free to use this work in any way permitted by the copyright and/or the licence that applies to your usage. For other uses, you must obtain permission from the rights-holders.

This document is made available under a Creative Commons license.



The license information is available online:

<https://creativecommons.org/licenses/by/4.0/legalcode>



OPEN ACCESS



Artificial Intelligence in Election Campaigns: Perceptions, Penalties, and Implications

Andreas Jungherr ^a, Adrian Rauchfleisch ^b, and Alexander Wuttke^c

^aInstitute of Political Science, University of Bamberg, Bamberg, Germany; ^bGraduate Institute of Journalism, National Taiwan University, Taipei City, Taiwan; ^cGeschwister-Scholl-Institute of Political Science, LMU, Munich, Germany

ABSTRACT

As political parties around the world experiment with Artificial Intelligence (AI) in election campaigns, concerns about deception and manipulation are rising. This article examines how the public reacts to different uses of AI in elections and the potential consequences for party evaluations and regulatory preferences. Across three preregistered studies with over 7600 American respondents, we identify three categories of AI use: campaign operations, voter outreach, and deception. While people generally dislike AI in campaigns, they are especially critical of deceptive uses, which they perceive as norm violations. However, parties engaging in AI-enabled deception face no significant drop in favorability, neither with supporters, opponents, nor independents. Instead, deceptive AI use increases public support for stricter AI regulation, including calls for an outright ban on AI development. These findings indicate that public disapproval of deceptive uses of AI does not directly translate into incentives for parties to forgo them, at least in the polarized political environment of the US.

KEYWORDS

Artificial intelligence; election campaigns; public opinion; political communication; regulatory governance

Artificial Intelligence (AI) is starting to feature in election campaigns around the world (Barredo-Ibáñez et al., 2021; Lin, 2023; Raj, 2024; Sifry, 2024; Swenson et al., 2024). From targeting voters to generating messages, AI offers political actors powerful new tools (Foos, 2024; Kruschinski et al., 2025; Tomić et al., 2023). Yet these technological advances are often met with sharp criticism from journalists and experts, who frequently highlight AI's perceived potential to fuel electoral deception and manipulation (e.g., Sanders & Schneier, 2024; Verma & Vynck, 2024). This tension between the opportunities AI offers to political actors and the public's growing concerns about its misuse raises urgent questions for democracies, particularly regarding the role of public opinion in constraining the behavior of political elites, shaping campaign tactics, and informing appropriate regulatory responses. This article contributes to that debate by providing evidence on how citizens perceive the use of AI in electoral campaigns and whether they sanction political actors seen as violating electoral norms.

CONTACT Andreas Jungherr  andreas.jungherr@gmail.com; andreas.jungherr@uni-bamberg.de  Institute of Political Science, University of Bamberg, Feldkirchenstraße 21, Bamberg 96052, Germany

 Supplemental data for this article can be accessed online at <https://doi.org/10.1080/10584609.2025.2611913>

Copyright © 2026 The Author(s). Published with license by Taylor & Francis Group, LLC.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited. The terms on which this article has been published allow the posting of the Accepted Manuscript in a repository by the author(s) or with their consent.

Across three preregistered studies with more than 7600 American respondents, we find that while the public strongly disapproves of AI-enabled deception in election campaigns, this disapproval does not translate into favorability penalties for those responsible. Instead, it leads to heightened support for strict AI regulation, including calls for a general ban on AI development. This reveals a misalignment between public worries about deceptive uses of AI and electoral incentives for parties to forgo them.

We introduce a new framework for understanding how people assess the use of AI in elections. We distinguish between three core types of AI use: improving campaign operations (e.g., automated content generation, chatbot-based communication with volunteers/supporters, and algorithmic segmentation of donor and walk lists), enhancing voter outreach (e.g., testing and optimizing message content and delivering personalized appeals across digital platforms), and engaging in deception (e.g., deepfakes). We argue that public reactions vary based on the perceived norm violation of each use, which in turn affects not only attitudes toward specific campaigns but also broader beliefs about democracy, personal control, and technology regulation.

Our findings are based on three studies: (1) a nationally representative survey of public attitudes toward various campaign-related uses of AI (Perceptions: $n = 1,199$), (2) a survey experiment testing how different AI use types affect political attitudes and regulatory preferences (Reactions: $n = 1,985$), and (3) a partisanship-based experiment assessing whether deceptive AI use leads to electoral penalties (Penalties: $n = 4,451$). Together, these studies provide a systematic account of how the public views AI in elections. This early evidence offers a baseline for the implementation and critical discussion of AI in political communication in future election cycles, taking into account different applications of AI, public attitudes, and the incentives of political actors.

Theoretical Framework: Public Reactions to AI Use in Elections

We propose a framework to explain how the public reacts to different uses of AI in election campaigns and how these reactions affect attitudes toward democracy and technology governance. Prior research shows that attitudes toward AI are not uniform. Instead, they seem to depend on how and to what purpose AI is used (Raviv, 2025; Zhang & Dafoe, 2019). Similarly, we do not expect public opinion on the uses of AI in election campaigns to be uniform. We expect people's reactions to depend on how AI is used and whether these uses violate norms of appropriate political conduct. These perceptions can influence judgments across multiple domains, including evaluations of campaigning practices, democratic legitimacy, personal autonomy, and support for regulation.

Our framework consists of four interconnected steps:

- (1) categorizing different types of AI use in campaigns;
- (2) identifying perceived norm violation as key mechanism;
- (3) outlining the key domains in which public attitudes are affected; and
- (4) anticipating the downstream consequences of norm-violating uses.

Step 1: Types of AI Use in Election Campaigns

A consistently growing set of academic studies and journalistic accounts from all over the world show that the use of AI in election campaigns is highly varied and extends far beyond the often-discussed threat of deepfakes (Barredo-Ibáñez et al., 2021; Dommett, 2023; Foos, 2024; Garimella & Chauchard, 2024; Kamal & Kaur, 2025; Kruschinski et al., 2025; Lin, 2023; Raj, 2024; Sifry, 2024; Swenson et al., 2024; Tomić et al., 2023). Being aware of this functional diversity is important in understanding the uses of AI for campaigning, given that campaigns typically select tools based on their anticipated contribution to specific campaign objectives and organizational needs (Earl & Kimport, 2011; Hindman, 2005; Jungherr et al., 2020).

To capture this functional diversity, we introduce an original categorization scheme that distinguishes three primary types of AI use in election campaigns: campaign operations, voter outreach, and deception. These categories emerged inductively through a review of documented AI applications in journalistic reporting and academic literature. The resulting framework reflects the distinct ways in which campaigns employ AI tools to pursue strategic goals as well as manage everyday operational demands.

AI in Campaign Operations

Campaigns use AI to streamline internal processes and improve organizational efficiency. Applications include automated content generation, chatbot-based communication with volunteers or supporters, and algorithmic segmentation of donor and walk lists (e.g., Carrasquillo, 2024; Chow, 2024; Foos, 2024; Sifry, 2024; Swenson, 2023). These uses tend to receive limited public attention but play a significant role in shaping how campaigns allocate resources, communicate internally, and adapt to fast-changing circumstances.

Past research on digital campaigning shows that technological adoption is often driven less by public-facing spectacle than by improvements to everyday processes (Hindman, 2005; Karpf, 2012; Kreiss, 2011, 2012; Nielsen, 2011). In this regard, AI continues the evolution of campaign infrastructure, offering new tools for scaling operations with limited staff and budgets.

AI in Voter Outreach

AI is also used to enhance voter outreach by identifying likely supporters, testing and optimizing message content, and delivering personalized appeals across digital platforms (e.g., Burke & Sunderman, 2024; Chatterjee, 2024; Hackenburg & Margetts, 2024). These practices build on a well-established tradition of data-driven campaigning, including microtargeting and experimental message testing (Dommett et al., 2024; Green & Gerber, 2004/2023; Hersh, 2015; McCarthy, 2020; Nickerson & Rogers, 2014; Turrow et al., 2012).

In this context, AI enables campaigns to refine their persuasive strategies and tailor messages to specific audiences more efficiently than before. While some outreach practices may raise concerns about manipulation or privacy, they are broadly consistent with conventional campaign goals and practices, which are now augmented by machine learning and generative tools.

AI in Deception

AI can also be used to intentionally mislead, impersonate, or obscure key information. Deceptive practices include the generation of synthetic audio or video that falsely portrays a candidate or opponent, the impersonation of political figures using deepfakes, and the deployment of AI-generated content in coordinated astroturfing campaigns (e.g., via bots or automated e-mail outreach to journalists and voters) (e.g., Barari et al., 2025; Bond, 2024; Coltin, 2024; Linvill & Warren, 2024; Ternovski et al., 2022).

Such uses build on familiar political tactics such as negative campaigning and strategic misinformation (Austen-Smith, 1992; Bucciol & Zarri, 2013; Haselmayer, 2019; Jay, 2010; Lau & Rovner, 2009; Nai, 2020). AI is feared to enhance these tactics by facilitating likeness, scalability, and speed, which may raise their potential impact. A Pew Research Center study from Summer 2024 found that Republicans and Democrats alike were equally worried about the impact of AI on the 2024 presidential campaign (Gracia, 2024). At the same time, early fears about widespread societal harm from AI-enabled deception, especially deepfakes, appear to be overstated, and evidence on their actual influence remains scarce (Simon et al., 2023).

We treat deception as a distinct category of AI use because it reflects an intent to influence public perception through misrepresentation, rather than merely an unintended side effect of technological deployment. As with operational and outreach applications, deceptive uses represent deliberate choices about how campaigns seek to achieve specific communication goals.

By distinguishing between operations, voter outreach, and deception, our framework offers a comprehensive account of how campaigns use AI, capturing applications that range from routine and arguably constructive to controversial and potentially harmful. Raising awareness of the different ways AI can be used in campaigns is especially important given that journalistic coverage focuses predominantly on one type of use: deception. A content analysis of news articles discussing AI in US elections ($n = 3,333$) shows that 63,58% addressed deceptive uses of AI. By contrast, only 8,58% discussed how AI might be used to improve campaign operations or for voter outreach (see Online Appendix 6 for details). As a result, the public portrayal of AI in political communication skews heavily toward deceptive uses, neglecting others.

Step 2: Norm Violation as Key Mechanism

In their choices of tactics, technologies, and communication strategies, campaigns are constrained by norms of acceptable political behavior and by the extent to which the public and press are willing to enforce these norms and sanction violations. This also applies to the use of AI in election campaigns.

Norms can be understood as shared, often implicit, expectations that define what is considered appropriate within a group or society (Cialdini & Trost, 1998). In interpersonal contexts, norms guide behavior and shape expectations. Violations of these norms often provoke negative reactions, including social sanctions or reduced credibility (Burgoon & Le Poire, 1993; van Kleef et al., 2015). This mechanism also extends to how people assess the actions of organizations and institutional actors (Dahl et al., 2003).

In electoral politics, actors who violate established norms often face public backlash. Such norm violations, such as excessive negativity or incivility, can reduce support for the

offending party or candidate, both in terms of attitude and vote choice (Ansolabehere et al., 1994; Fridkin & Kenney, 2011; Muddiman, 2017). We argue that the same evaluative logic applies to the use of AI: people are likely to judge parties' AI uses based on their perceived adherence to, or deviation from, normative expectations of legitimate political competition.

Crucially, not all uses of AI are equally likely to be seen as norm violations. Deceptive uses, such as generating fake content or impersonating political figures, clearly conflict with general social norms. In many domains, deception leads to negative affective and behavioral responses toward the deceiver (Croson et al., 2003; Ohtsubo et al., 2010; Tyler et al., 2006). We therefore expect that when people learn a party has used AI deceptively, they will react with disapproval, perceiving the act as a violation of democratic norms.

The case of voter outreach is more ambiguous. On the one hand, critics have long argued that targeting voters with individually tailored messages, especially based on psychological profiling, threatens informational autonomy and undermines democratic deliberation (Bayer, 2020; Bennett & Manheim, 2006). Survey evidence also shows that people often feel uneasy about behavioral targeting by both commercial and political actors (McCarthy, 2020; Turrow et al., 2012). On the other hand, such techniques may now be sufficiently normalized that people no longer view them as norm violations, even if they remain uncomfortable with them. Moreover, targeted outreach may help parties reach disengaged or hard-to-reach citizens in fragmented media environments (Dommett et al., 2024; Jungherr et al., 2020). As such, public reactions to AI-enabled targeting may depend on whether it is framed as manipulative or as a means of broadening participation.

Finally, the use of AI to improve internal campaign operations, such as scheduling, donor segmentation, or content drafting, is unlikely to be viewed as normatively problematic, given the mundane nature of these uses.

In sum, we expect perceived norm violation to be a key mechanism shaping public responses to campaign uses of AI. The intensity and direction of those responses, however, are likely to vary by context and function: strongly negative in the case of deception, more ambivalent in the case of voter outreach, and potentially positive in the case of operational uses.

Step 3: Related Attitudes

Public experiences with AI use in politics are likely to shape more than just opinions about election campaigns. Because campaigns attract intense media coverage and public scrutiny, they often serve as high-profile contexts in which emerging technologies are introduced, contested, and symbolically evaluated. As such, campaign-related AI uses frequently become exemplars (Murphy, 2002), cases that people refer to when forming judgments about AI's broader role in politics and society, especially in areas where personal experience or direct evidence is limited.

Historical cases illustrate this dynamic. For example, the campaigns of Howard Dean and Barack Obama became emblematic of the democratic potential of digital technologies, contributing to hopeful narratives about online mobilization and participatory empowerment (Kreiss, 2016; Shirky, 2008). In contrast, accounts of Cambridge Analytica's role in the Brexit and Trump campaigns have become touchstones in narratives about the dangers of data-driven surveillance and political manipulation, fueling public demands for stronger oversight of digital platforms (Weiss-Blatt, 2021; West & Allen, 2020).

Likewise, contemporary uses of AI in campaigns, whether celebrated or condemned, may shape public attitudes across a range of domains. At the most immediate level, individuals form judgments about specific campaign applications of AI: whether they are appropriate, effective, or concerning. These evaluations can then extend to more general political beliefs, such as perceptions of party integrity, election fairness, or the overall quality of democracy. Over time, campaign-related exemplars may even influence how people perceive AI's role in their personal lives, particularly in relation to autonomy and control. For example, encountering deceptive AI uses in politics may heighten feelings of disempowerment or suspicion toward algorithmic systems more broadly.

Finally, these perceptions can also shape public attitudes toward AI governance. If political uses of AI are seen as norm-violating or harmful, they may drive support for more restrictive or precautionary regulation, not just within the political sphere but across all domains of AI development and deployment. While public support of regulation does not automatically drive regulation, it is indicative of how the public sees technology and its risks and opportunities. In this way, campaign experiences with AI can have ripple effects, shaping how societies understand emerging technologies well beyond the electoral arena.

Step 4: Anticipating Downstream Effects

Building on the previous steps, we can now specify our expectations regarding how different uses of AI in election campaigns affect public attitudes. When campaign uses of AI are perceived as norm violations, we expect a cluster of negative reactions. These include:

- Lower favorability toward the implicated party;
- Diminished trust in democratic institutions, including the fairness of elections;
- Reduced sense of personal autonomy and control; and
- Stronger support for regulatory oversight, including more restrictive or precautionary AI policies.

These reactions are not limited to evaluations of campaign tactics; they extend to broader judgments about democracy, technology, and the legitimacy of emerging systems of governance. In this sense, norm-violating AI uses in politics can act as focal points that shape public opinion across multiple domains of political and personal life.

However, because these responses occur in a political context, they are unlikely to be uniform across the electorate. A robust body of research shows that political partisans are prone to motivated reasoning: they tend to discount or rationalize information that reflects poorly on their preferred party, while reacting strongly to norm violations by political opponents (Jost et al., 2013; Kahan, 2016; Williams, 2023). We therefore expect heterogeneous effects across partisan lines. While norm-violating AI use may provoke strong disapproval in general, this disapproval is likely to be asymmetrical, more intense among opponents of the implicated party and independents, while remaining more muted or absent among its supporters.

Observational Survey (n = 1,199)

Quota Sampling: Age, gender, region, & education
 Representative for American population 18+
 Provider: Ipsos
 Field Time: April 4 - April 17, 2024

**Survey Experiment (n = 1,985)**

Stratified Sampling: Age, sex, & ethnicity
 Representative for American population 18+
 T1: Treatment, AI for Deception (n = 497)
 T2: AI for Campaign Operations (n = 494)
 T3: AI for Voter Outreach (n = 497)
 C: Control, No Information (n = 497)
 Provider: Prolific
 Field Time: June 19 and June 21, 2024

Survey Experiment (n = 4,451)

Sample: Republican Partisans (n = 1,485)
 Sample: Democrat Partisans (n = 1,489)
 Sample: Independents (n = 1,477)
 T1: Deceptive Uses of AI by Democrats
 T2: Deceptive Uses of AI by Republicans
 C: Control, No Information
 Provider: Prolific
 Field Time: June 25 and June 30, 2024



Figure 1. Research design.

This dynamic helps explain a core tension explored in this study: although people express strong negative views about deceptive AI use, these attitudes do not always translate into electoral penalties for the responsible parties. Instead, motivated reasoning may buffer in-group parties from reputational harm, even as public concern drives broader demands for regulatory intervention.

Materials and Methods

To test our framework for the explanation of public reactions to AI uses in election campaigns, we ran three preregistered surveys in the US, including two survey experiments (for research design see [Figure 1](#) and Online Appendix). All studies were approved by the Institutional Review Boards at the University of Bamberg and the National Taiwan University. Each study focuses on a distinct aspect of the framework: public attitudes toward different AI uses (Study 1: Perceptions), causal effects of exposure to those uses (Study 2: Reactions), and whether parties are penalized for using deceptive AI (Study 3: Penalties). With this study, we present the first systematic evidence on how Americans think about different uses of AI in election campaigns. We choose the United States as a case to study public attitudes toward AI in campaigning, since American campaigns have historically been early adopters of new technologies, and AI already figures prominently in U.S. journalism and policy debates. This makes the U.S. a promising context for both the use of AI in campaigning activities and heightened public awareness of AI. At the same time, the U.S. remains a special case, particularly due to its strong partisan

polarization, which may lead to partisan-specific patterns that do not automatically translate to other political contexts. Nevertheless, given the AI- and campaigning-specific dynamics at play, the U.S. provides a valuable environment to examine early public attitudes toward the use of AI in democratic campaigning.

Study 1 (Perceptions): Observational Survey (n = 1,199)

In Study 1, we queried people for their opinions on specific uses of AI in elections and their views on the benefits and risks of AI in other areas. We ran a preregistered survey ($n = 1,199$) among members of an online panel that the market and public opinion research company Ipsos provided. We used quotas on age, gender, region, and education to realize a sample representative of the US electorate (See Online Appendix 1 for details). The survey was fielded between April 4 and April 7, 2024. Before running the survey, we registered our research design, analysis plan, and hypotheses about outcomes.¹ We did not deviate from the registered procedure. We provided respondents with short descriptions of various campaigning tasks for which parties and candidates use AI. These tasks fall into three broad categories: (I) campaign operations, (II) voter outreach, (III) deception. For each category, we identified five example tasks that practitioners and journalists have documented and discussed. We showed each respondent three randomly drawn tasks for each category and asked them whether this use of AI in elections: (1) worried them, (2) felt like a norm violation, (3) was likely to make politics more interesting to voters, and (4) increase participation. We combined items “This use of AI makes politics more interesting to voters,” and “This use of AI increases voter engagement” into one index (Rise in Voter Involvement) to capture AI’s likely impact on voter interest and mobilization. As two distinct independent variables, we also measured AI risk ($\alpha = 0.78$) and benefit perceptions ($\alpha = 0.84$) with three items each. For each dependent variable, we estimated a multilevel model with varying intercepts for cases (15) and participants. Furthermore, as specified in the preregistration, we used data imputation to fill in the missing responses for AI benefits and AI risks (both models, with and without data imputation, show the same results; see Online Appendix 5.1). See Online Appendix 1.1 for an overview of variables, question wordings, operationalizations, key diagnostics of item measurements, and the data imputation procedure for Study 1.

Study 2 (Reactions): Survey Experiment (n = 1,985)

In our second study, we test the causal effects of learning about different types of AI use in elections. We ran a preregistered survey experiment with members of an online panel provided by Prolific ($n = 1,985$).² We used quotas to realize a sample resembling the US electorate (See Online Appendix 1.2 for details). The survey was fielded between June 19 and June 21, 2024. We divided respondents into three treatment groups and one control group (C, $n = 497$). Treatment 1 (T1, $n = 497$) contained information about campaigns’ uses of AI for deception. Treatment 2 (T2, $n = 494$) contained information about campaigns’ uses of AI for improving campaign operations. Treatment 3 (T3, $n = 497$) contained information about campaigns’ uses of AI for voter outreach. The dependent variables in this study cover four broad areas: campaign use (worry, norm violations, and positive impact on politics), democracy (fairness of elections, favorability of parties in general and specific parties), lifeworld (personal loss of control), and regulation (of AI use in elections, AI in general, priority of regulation over innovation, and support for an AI moratorium). We worded the moratorium item to mirror concerns from

a widely discussed public call³ to halt AI development, anchoring the measure in real-world discourse (see Online Appendix 1.2). We preregistered our research design, analysis plan, and hypotheses about outcomes before the survey.⁴ We did not deviate from the preregistered procedure. For detailed information on treatments, variables, question wordings, operationalizations, key diagnostics of item measurements, and manipulation checks for Study 2, see Table 11 in the Online Appendix.

Study 3 (Penalties): Survey Experiment ($n = 4,451$)

In Study 3, we test whether parties face a penalty for deceptive AI uses attributed to them and whether partisans' group-protective cognitions lead to heterogeneous effects of being informed about deceptive uses. For this preregistered study,⁵ we recruited three samples containing only respondents identifying as partisans for (1) Democrats ($n = 1,489$), (2) Republicans ($n = 1,485$), or as (3) Independent ($n = 1,477$). Prolific prescreened partisans. No attempt to be representative was made. The survey was fielded between June 25 and June 30, 2024. Respondents in these three samples were exposed to either of two treatments or were assigned a pure control group that did not receive any information. Treatments contained information about deceptive uses of AI by candidates from the Democratic Party (T1) or the Republican Party (T2). This allows us to identify whether group-protective cognition leads partisans to discount information about uses of AI by parties they support, compared to adjusting related attitudes when being informed about deceptive uses by parties they oppose, and how this compares to reactions by Independents. For the dependent variables, we focused on three broad areas: campaign use (worry, norm violations, and positive impact on politics), democracy (fairness of elections, favorability of specific parties), and regulation (priority of regulation over innovation and support for an AI moratorium). We preregistered our research design, analysis plan, and hypotheses about outcomes before the survey. We did not deviate from the preregistered procedure. For detailed information on treatments, variables, question wordings, operationalizations, key diagnostics of item measurements, and manipulation checks for Study 3, see Online Appendix 1.3.

Results⁶

Study 1, Perceptions: Public Disapproval of AI in Campaigns Varies by Use

We asked a representative sample of Americans ($n = 1,199$) for their opinions on specific uses of AI (see Online Appendix 1.1 for details). In our preregistered study, we provided respondents with a selection from fifteen short descriptions of various campaigning tasks for which AI might be used. These tasks fall into three functional categories:

- Campaign operations: including automated content generation, chatbot-based communication, and the segmentation of donor and walk lists.
- Voter outreach: including identifying persuadable voters, optimizing message appeal (either at scale or individually), and generating personalized ads.
- Deception: including the undeclared use of AI to create misleading media (e.g., deepfakes), impersonate candidates, or conduct coordinated astroturfing via bots and large language models.

Table 1. Share responses that agree with assessment (Study 1).

Campaign Task	Worry (in %)	Norm Violation (in %)	Rise Voter Involvement (in %)
Deception: Astroturfing (interactive)	68.57	64.49	33.88
Deception: Astroturfing (social media)	76.37	69.86	32.18
Deception: Deceptive Robocalls	68.78	64.14	28.69
Deception: Deepfakes (negative campaigning)	71.58	68.26	29.46
Deception: Deepfakes (self-promotional)	71.80	68.76	36.23
Outreach: Ad Optimization	63.54	57.78	39.02
Outreach: Data Driven Targeting	56.11	55.90	34.16
Outreach: Fundraising	58.41	55.09	39.38
Outreach: Message Testing & Opinion Research	60.83	54.92	35.63
Outreach: Outreach Optimization	60.91	58.85	32.30
Operations: Automating Interactions	67.72	63.56	35.84
Operations: Deepfakes (benign)	56.43	57.11	40.18
Operations: Resource Allocation	49.80	47.83	41.70
Operations: Transcription	53.83	53.42	33.13
Operations: Writing	62.91	55.10	29.50

Each category included five representative tasks. For each task they were shown, respondents were asked how much the use worried them, whether they perceived it as violating campaign norms, and whether they believed it could increase voter engagement.

Table 1 summarizes the proportion of respondents who rated each item above the midpoint on a 7-point scale (responses > 4), excluding missing values. The results show clear distinctions across categories. Respondents were most concerned about deceptive uses of AI, which they saw as more likely to violate democratic norms and less likely to boost voter involvement. While all AI uses met with some degree of concern, certain campaign operations, particularly low-profile, non-communicative tasks like resource allocation and transcription, were viewed most favorably. Outreach-related uses were generally seen as more problematic than these operational tasks, but less troubling than deceptive practices. However, operational tasks involving AI-generated interaction or writing also drew substantial concern, underscoring that reactions vary within categories as well as across them.

Figure 2 visualizes the distribution of responses across the fifteen use cases by category. The ridgeline plots show that, although public opinion on AI in politics is generally cautious, deceptive uses consistently meet the strongest disapproval. Across all three measures (i.e. worry, norm violation, and perceived potential to increase engagement), AI for deception is judged most negatively. These plots also highlight the internal variation within each category, reinforcing the importance of distinguishing between specific applications of AI in campaign contexts.

To further examine these differences, we estimated regression models predicting levels of worry, perceived norm violation, and belief in increased voter involvement ($n = 1,199$). As shown in the top row of Figure 3, even after controlling for individual characteristics, deceptive uses of AI consistently met with greater concern and were more likely to be viewed as norm violations, while also being least likely to be seen as enhancing democratic participation. Estimates for operational and outreach uses overlap, suggesting that respondents perceive these uses in broadly similar ways but distinct from how they view deceptive applications.

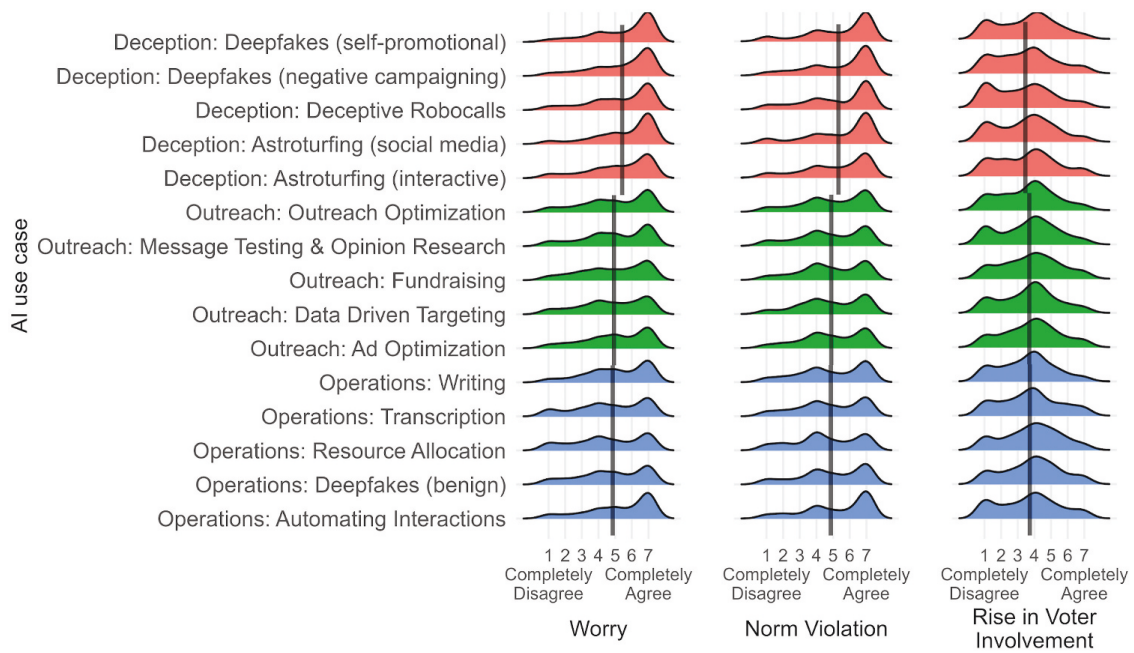


Figure 2. Attitudes toward AI uses in elections by type. The vertical line indicates the mean per category. The deception category differs significantly ($p < .001$) from the other two for all three outcome variables.

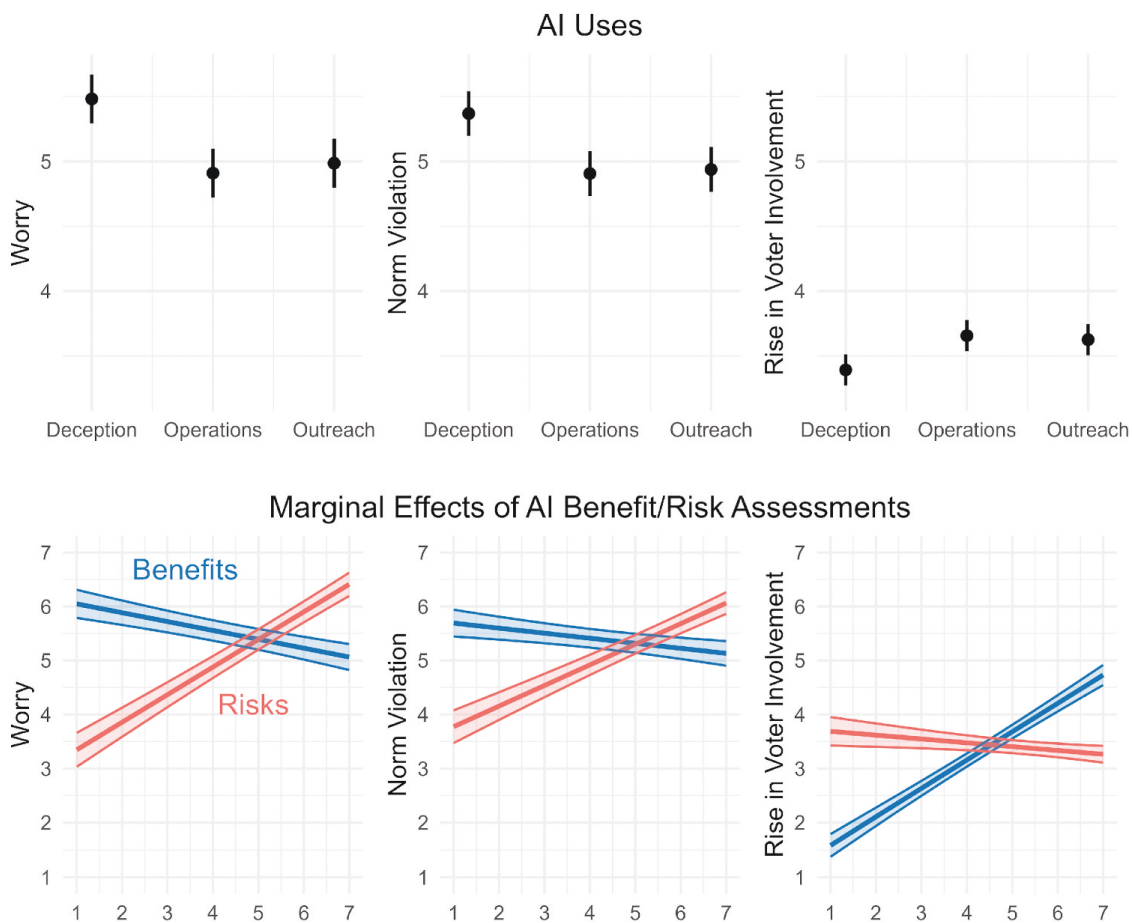


Figure 3. Attitudes toward AI uses in elections, regressions (for campaign tasks, deception is used as a reference group). Estimates with 95%-CIs. The deception category differs significantly ($p < .001$) from the other two for all three outcome variables. Also, both AI benefit and risk perception are significant for all three outcome variables.

Study 2: Reactions to AI Uses

In a preregistered follow-up experiment ($n = 1,985$), we examined the causal effects of learning about different types of AI use in election campaigns (see Online Appendix 1.2 for details). Respondents were randomly assigned to one of three treatment groups or a control group ($n = 497$). Each treatment provided information about a different category of campaign AI use:

- Deception Treatment ($n = 497$): presented examples of AI-enabled deception by campaigns.
- Operations Treatment ($n = 494$): presented examples of AI used to improve internal campaign processes.
- Outreach Treatment ($n = 497$): presented examples of AI used to enhance voter targeting and communication.

Since Study 1 showed that deceptive uses of AI were perceived most negatively, we preregistered the deception condition as the reference group.

The results indicate that perceived norm violation plays a central role in shaping public reactions, but only for specific outcome dimensions (see Figure 4). Respondents who received the deception treatment expressed significantly higher levels of worry and were more likely to view the behavior as a violation of campaign norms, compared to all other groups.

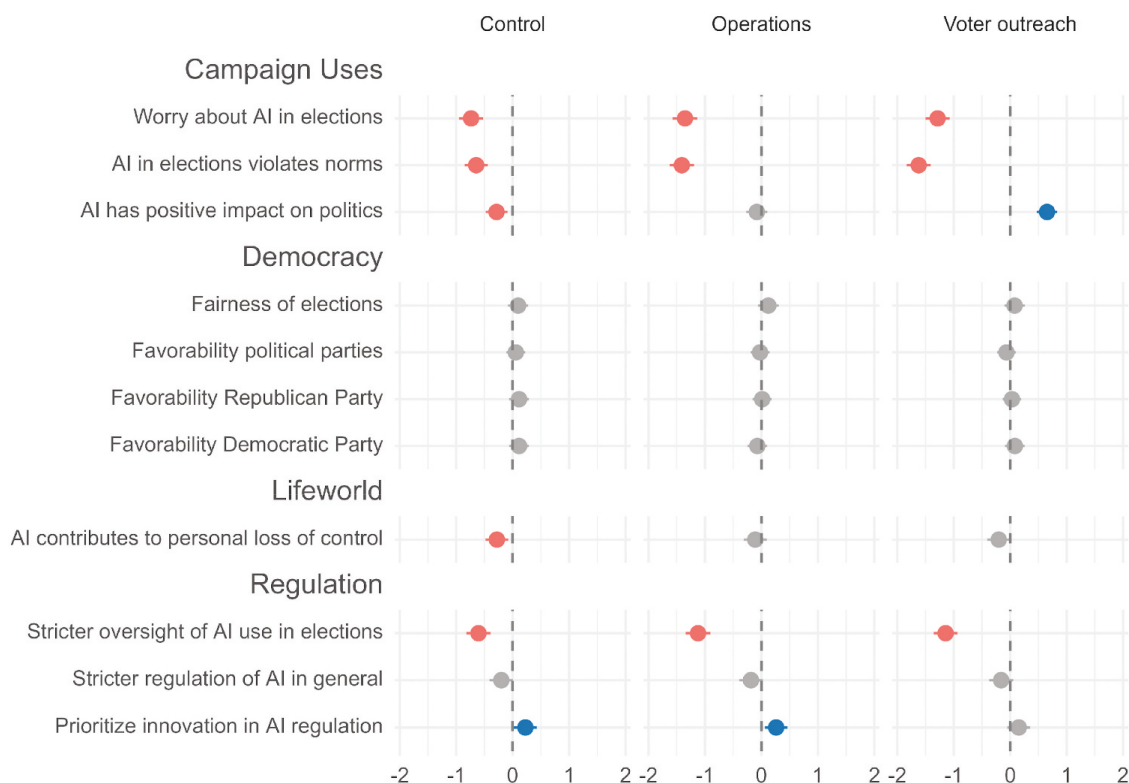


Figure 4. Effects of information about different uses of AI in elections (reference category: deception). Estimates with 95%-CIs. Grey estimates are not significant.

However, the effects did not extend to broader democratic attitudes. Exposure to AI-enabled deception did not significantly affect perceptions of electoral fairness or party favorability. In other words, while deceptive AI use elicits disapproval, it does not appear to carry an electoral penalty, at least not in terms of reduced support for the party involved or diminished trust in the election process. Accordingly, we tested in a follow-up study directly whether parties suffer a favorability penalty when deceptive AI use is attributed to them.

Study 3: Parties face no favorability penalty for deceptive AI use

This preregistered study used three samples composed exclusively of self-identified partisans: Democrats ($n = 1,489$), Republicans ($n = 1,485$), and Independents ($n = 1,477$) (see Online Appendix 1.3 for details). Within each group, respondents were randomly assigned to one of two treatment conditions or a control group. The treatments provided information about deceptive AI use by either a Democratic or a Republican candidate.

The results provide no evidence that deceptive AI use reduces support for the implicated party (see Figure 5). While both Democratic and Republican respondents expressed a greater sense of norm violation when exposed to information about their own party's

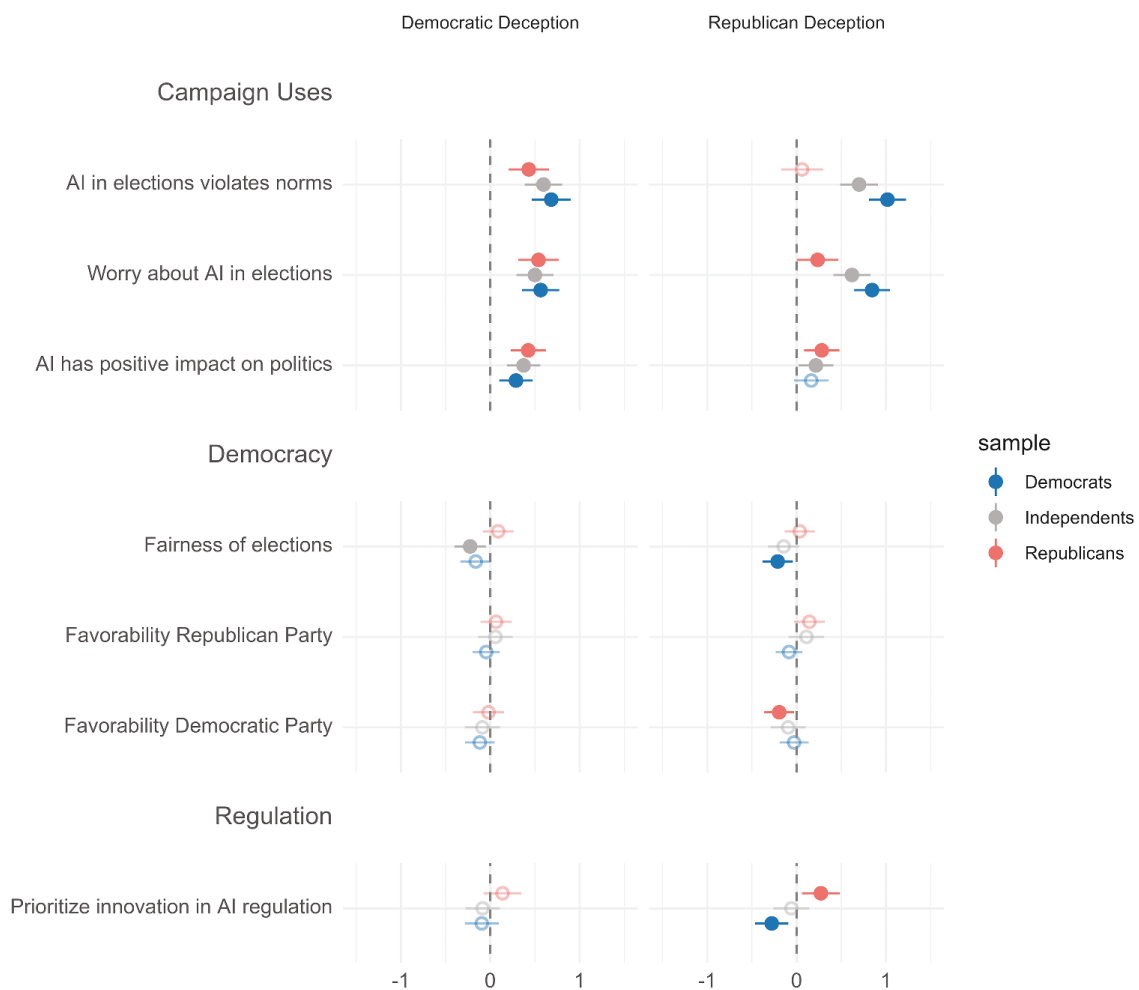


Figure 5. Effects of information about alleged deceptive uses of AI for partisans and Independents (reference category: control). Estimates with 95% CIs. Hollow estimates are not significant.

alleged misconduct, they did not lower their favorability ratings compared to the control group. In other words, partisans disapprove of deceptive AI use in principle but are unwilling to penalize their own party for it in practice.

This pattern extends to Independents, who also did not significantly change their favorability ratings in response to either treatment condition. Equivalence tests show that no substantial shifts in party favorability occurred among any group (see Online Appendix 4).

Together with the findings from Study 2, this result highlights a critical disconnect: although norm-violating uses of AI generate disapproval, they do not translate into political consequences for the parties responsible. This suggests that the public's normative expectations and attitudinal reactions are not sufficient to constrain campaign behavior, especially in an environment shaped by partisan loyalty and motivated reasoning.

Studies 2 and 3: AI-Enabled Deception Increases Support for Restrictive AI Governance

Beyond attitudes directly related to campaigning, Studies 2 and 3 also examined how exposure to AI-enabled deception influences broader views, particularly assessments of personal autonomy and support for AI regulation. These outcomes extend the analysis beyond electoral dynamics to public perceptions of AI's role in society and the legitimacy of its continued development.

As shown in [Figure 4](#), respondents in Study 2 who were informed about deceptive campaign uses of AI were significantly more likely to report a sense of diminished personal control. This suggests that people use high-salience political examples, such as deceptive AI in election campaigns, as heuristics for evaluating the impact of AI on their daily lives and civic agency.

The effect on regulatory attitudes was even more pronounced. Support for a complete stop to AI development and use was significantly higher among those exposed to information about AI deception in elections. While 29% of respondents in the control group supported such a ban, 38% of those in the deception treatment group did so. These respondents also expressed stronger support for stricter oversight of AI use in elections and favored prioritizing safety over innovation in the regulation of emerging technologies (see [Figure 6](#)).

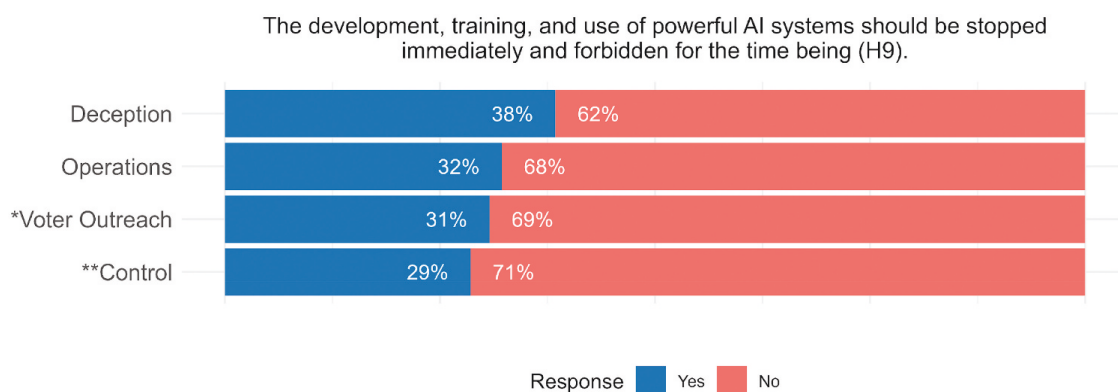


Figure 6. Effects of information about different uses of AI in elections on governance preferences (reference category: deception), $p < .05$ (*), $p < .01$ (**), $p < .001$ (***).

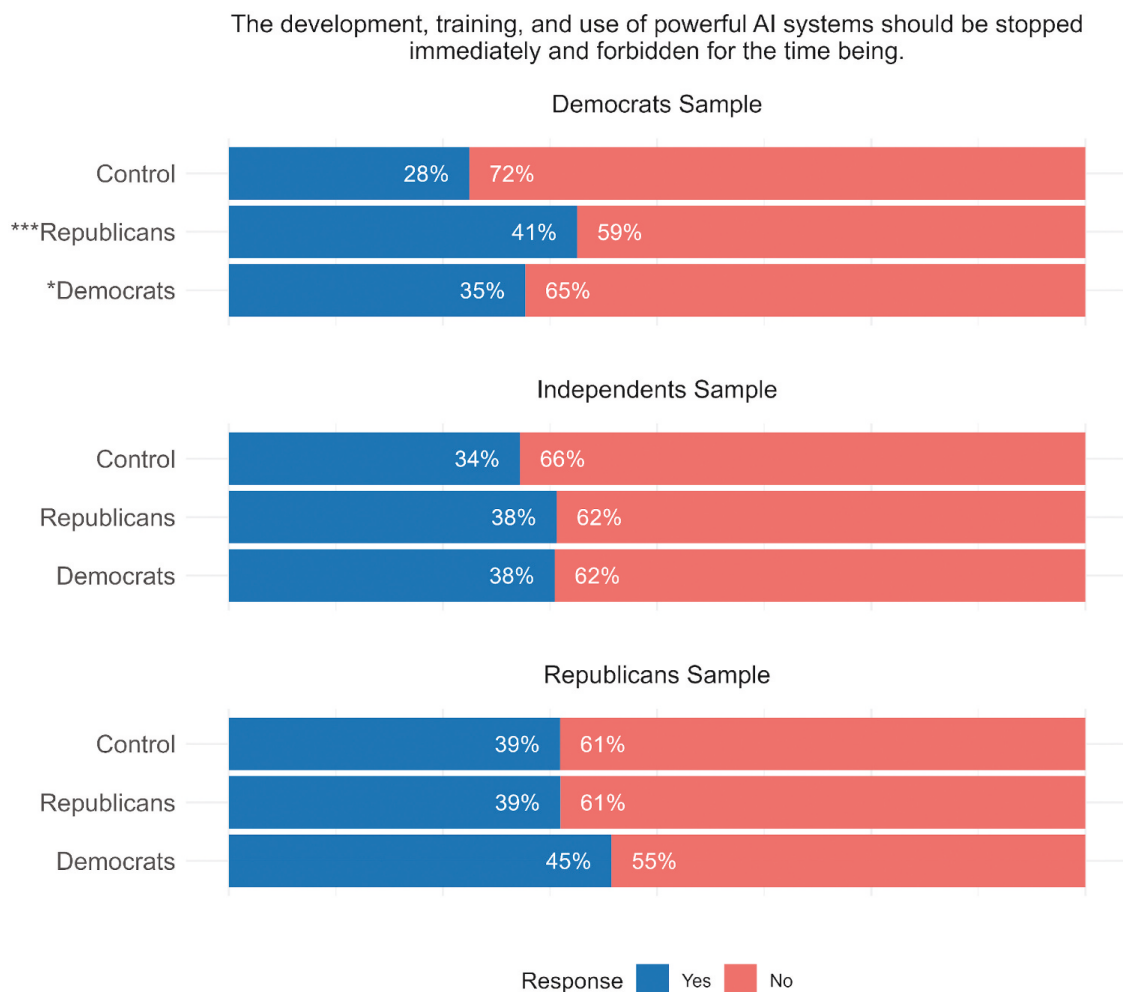


Figure 7. Effects of information about alleged deceptive uses of AI on governance preferences for partisans and independents (reference category: control), $p < .05$ (*), $p < .01$ (**), $p < .001$ (***)

The findings from Study 3 corroborate these results. Again, information about AI deception increased support for halting AI development (see Figure 7). While baseline support for an AI ban varied by partisanship, 39% among Republicans and 28% among Democrats in the control group, Democrats significantly increased their support when exposed to deceptive AI use, regardless of whether the behavior was attributed to their own party or the opposing one.

These results highlight a critical asymmetry in the political consequences of AI deception. Although parties do not face direct electoral penalties for deceptive AI use (as shown in Study 3), the practice generates negative externalities: it erodes public trust, undermines perceived autonomy, and increases support for more restrictive, and potentially hostile, approaches to AI governance. Notably, we find that exposure to deceptive AI increases public support for an outright halt to AI development. While this preference for safety-first regulation is rational considering perceived threats, it may carry unintended consequences. Calls for sweeping bans or highly precautionary regimes can hinder experimentation, limit adaptive learning, and delay the development of non-deceptive applications. Public concern, when expressed through blunt regulatory preferences, may inadvertently produce policy that slows innovation without significantly reducing risk. In this way, political uses of AI serve as symbolic exemplars of the technology's broader societal risks, helping to shape the trajectory of regulation and public acceptance far beyond the electoral domain.

Discussion

This study provides a systematic account of how people perceive and react to the use of AI in election campaigns. Across three preregistered studies, we find that the public broadly disapproves of AI in politics, especially when it is used deceptively. But this disapproval does not translate into meaningful electoral penalties for the parties involved. Instead, public exposure to deceptive AI practices in campaigns leads to downstream effects that extend beyond politics: heightened support for restrictive AI regulation, increased feelings of personal disempowerment, and stronger preference for safety over innovation in AI governance.

These findings support our theoretical framework. First, the public differentiates clearly between types of AI use: campaign operations and voter outreach are viewed more neutrally, while deception is consistently seen as a norm violation. Second, AI uses that are seen as norm violation also met with negative reactions. Third, these reactions are not confined to campaign evaluations; they spill over into broader domains, including personal autonomy, and regulatory preference. Finally, while people disapprove of deceptive uses, they do not penalize the parties responsible, suggesting a misalignment between public values and political incentives.

This misalignment matters. Political actors may gain competitive advantages from using AI deceptively while facing little or no electoral cost. Our experiments show that exposure to such uses can produce externalities by increasing public support for more restrictive AI regulation, including moratoria. Clearly, heightened support does not automatically translate into policy change. In the U.S., salient issues such as gun control illustrate how strong and persistent public concern can coexist with legislative inaction. We therefore interpret these shifts as symbolic pressures on the policymaking environment rather than evidence of imminent responsiveness. Future research should specify when and under which conditions public opinion on AI shapes agenda setting, coalition formation, and regulatory outcomes.

Not all uses of AI in elections have these implications. Our findings show, the public evaluates AI tools used to improve operations or voter outreach with greater nuance than deception. This underscores the importance of analyzing the full spectrum of electoral AI uses rather than focusing solely on AI-enabled deception (Foos, 2024; Kruschinski et al., 2025; Tomić et al., 2023), particularly as AI becomes more visible in democratic processes and the public arena (Jungherr, 2023; Jungherr & Schroeder, 2023).

Our results also underscore the symbolic power of electoral AI use. Because election campaigns are high-visibility events, they shape how people think about AI more generally. Deceptive uses in politics can become exemplars of technological risk, leading to reactions that affect AI governance in domains far beyond elections. This amplifies the responsibility of political actors, regulators, and professional observers like journalists and academics alike.

That responsibility is heightened by a shift we do not analyze empirically here: AI lowers production costs, enabling supporters – not only official campaigns – to create near-professional content at scale. Such grassroots uses may increasingly shape campaign narratives and could matter more than the public-facing AI deployments of party organizations. For parties, this raises questions of coordination, brand control, and accountability; for observers, it raises the need to assess provenance and degrees of affiliation between content and campaigns. These dynamics ask for systematic study in future research.

The framework presented here (i.e. the differentiation of AI use types, the centrality of perceived norm violations, and the mapping of downstream effects) provides a generalizable structure for cross-national and longitudinal work. While our data come from the U.S. electorate, future studies can apply this framework to other political systems. In countries with multi-party dynamics, stronger privacy norms, different media ecologies, or less polarized electorates, both public reactions and campaign behavior may differ in important ways (Dommett et al., 2024). For instance, outreach-related AI uses may trigger stronger reactions in contexts where behavioral targeting is culturally or legally contested, like Europe. Or deceptive practices might lead to less of a backlash in more permissible campaign contexts, like Asian democracies.

In addition, building on prior research into varying patterns of technological adoption and the incentives driving them (Bimber et al., 2012; Earl & Kimport, 2011; Jungherr et al., 2019) future studies should investigate how political actors evaluate the risks and benefits of adopting AI. How do resource-rich versus resource-constrained campaigns decide whether to use AI tools? Are some types of parties (e.g., system challengers or populists) more willing to engage in deceptive uses? Qualitative studies of campaign staff and consultants could shed light on these organizational dynamics and further inform regulatory design.

Methodologically, the survey and experimental tools developed here can be adapted for comparative studies or used to track change over time. As AI tools and public familiarity evolve, longitudinal designs will be essential to capturing shifting baselines of acceptance or concern.

Despite the strengths of our design, limitations remain. Our treatments were short textual descriptions. This may underestimate the real-world effects of repeated exposure to emotionally charged or visually engaging content, as well as follow-up coverage in the media and conversations in social circles. Also, the strong partisan divides in the U.S. may limit generalizability to less polarized democracies. Comparative and qualitative work is therefore essential to contextualizing these findings and extending them to diverse democratic environments. Finally, our study captures only isolated instances of exposure to information about deceptive uses of AI. While we did not observe immediate declines in perceptions of electoral fairness or party favorability, the absence of short-term effects does not preclude longer-term risks. Repeated norm-violating behavior without political consequences may gradually undermine citizens' confidence that elections are conducted on a level playing field. Over time, such erosion of trust in electoral integrity could itself constitute a democratic harm. Future research should therefore examine the cumulative effects of repeated exposure and the potential for longer-term consequences.

How AI is used in election campaigns and which of these uses are highlighted in public and academic discourse matters not only for political outcomes, but also for democratic norms, public trust, and the broader trajectory of technological governance. If deceptive applications come to dominate the public imagination of AI in politics, we risk not only undermining confidence in the integrity of electoral competition but also eroding trust in AI innovation more broadly. To counter this, public, academic, and regulatory efforts must capture the full range of electoral AI use cases. How citizens come to view both elections and AI will depend on this broader, more nuanced understanding.

Notes

1. Prereg Study 1: <https://osf.io/3nrb4>.
2. Prior research shows that Prolific provides excellent data quality. Peer et al. (2022), for example, show that Prolific provides even higher data quality than widely used panels such as Qualtrics and Dynata. For Study 1, which focuses on general attitudes toward different AI use categories, we relied on a high-quality Ipsos panel. In Studies 2 and 3, which emphasize experimental effects, we used Prolific with quota sampling, as detailed in the Online Appendix. Given the sampling approach and the large sample sizes (which are especially important for statistical power in experimental designs), we are confident in the generalizability of our findings.
3. <https://moratorium.ai>.
4. Prereg Study 2: <https://osf.io/wsrkv>.
5. Prereg Study 3: https://osf.io/vugp8/?view_only=5e4387422dc94458bb355e6e2e5fba3d.
6. For an overview of all preregistered hypotheses with all estimates and statistical tests, see Appendix 2, and for the complete tables of the models, see Appendix 5.

Acknowledgments

A. Jungherr, A. Rauchfleisch, and A. Wuttke contributed equally to the project. We used ChatGPT 4o for language improvement and editing. Correspondence and requests for materials should be addressed to A. Jungherr.

Disclosure Statement

No potential conflict of interest was reported by the author(s).

Funding

A. Jungherr's work was supported by a grant from the Volkswagen Stiftung and subsequently by a grant by the Bavarian State Ministry of Science and the Arts coordinated by the Bavarian Research Institute for Digital Transformation (bidt). A. Rauchfleisch's work was supported by the National Science and Technology Council, Taiwan (R.O.C.) (grant no. 113-2628-H-002-018 and 114-2628-H-002-007) and by the Taiwan Social Resilience Research Center (grant no. 114L9003) from the Higher Education Sprout Project by the Ministry of Education in Taiwan.

Notes on contributors

Andreas Jungherr holds the *Chair for Political Science, especially Digital Transformation* at the University of Bamberg and is Director at the *Bavarian Research Institute for Digital Transformation (bidt)*. He examines the impact of digital media on politics and society, with a special focus on artificial intelligence, political communication, and governance. He is the author of *Retooling Politics: How Digital Media is Shaping Democracy* (with Gonzalo Rivero and Daniel Gayo-Avello, Cambridge University Press: 2020) and *Digital Transformations of the Public Arena* (with Ralph Schroeder, Cambridge University Press: 2022).

Adrian Rauchfleisch is a Professor at the Graduate Institute of Journalism, National Taiwan University. His research focuses on the interplay of politics, technology, and journalism in Asia, Europe, and the United States. His new project explores Artificial Intelligence's influence on society across different cultural contexts.

Alexander Wuttke investigates the state of liberal democracy from the perspective of ordinary citizens and it explores the challenges and opportunities afforded by digitalization. He is an assistant professor in digitalization and political behavior at LMU Munich.

ORCID

Andreas Jungherr  <http://orcid.org/0000-0003-2598-2453>

Adrian Rauchfleisch  <http://orcid.org/0000-0003-1232-083X>

Data Availability Statement

Preregistrations, data, and analysis scripts are available at the project's OSF repository:

Study 1: Perceptions

•Prereg: <https://osf.io/3nrb4>

•Data and code: <https://osf.io/gheqz>

Study 2: Reactions

•Prereg: <https://osf.io/wsrkv>

•Data and code: <https://osf.io/8s7ye>

Study 3: Penalties

•Prereg: https://osf.io/vugp8/?view_only=5e4387422dc94458bb355e6e2e5fba3d

•Data and code: <https://osf.io/r3qa4>

Open scholarship



This article has earned the Center for Open Science badges for Open Data, Open Materials and Preregistered. The data and materials are openly accessible at <https://osf.io/gheqz> (Study 1), <https://osf.io/8s7ye> (Study 2), and <https://osf.io/r3qa4> (Study 3).

References

- Ansolabehere, S., Iyengar, S., Simon, A., & Valentino, N. (1994). Does attack advertising demobilize the electorate? *American Political Science Review*, 88(4), 829–838. <https://doi.org/10.2307/2082710>
- Austen-Smith, D. (1992). Strategic models of talk in political decision making. *International Political Science Review*, 13(1), 45–58. <https://doi.org/10.1177/019251219201300104>
- Barari, S., Lucas, C., & Munger, K. (2025). Political deepfakes are as credible as other fake media and (sometimes) real media. *The Journal of Politics*, 87(2), 510–526. <https://doi.org/10.1086/732990>
- Barredo-Ibáñez, D., De la-Garza-Montemayor, D.-J., Torres-Toukoumidis, Á., & López-López, P.-C. (2021). Artificial intelligence, communication, and democracy in Latin America: A review of the cases of Colombia, Ecuador, and Mexico. *Profesional de La Información*, 30(6), e300616. <https://doi.org/10.3145/epi.2021.nov.16>
- Bayer, J. (2020). Double harm to voters: Data-driven micro-targeting and democratic public discourse. *Internet Policy Review*, 9(1), 1–17. <https://doi.org/10.14763/2020.1.1460>
- Bennett, W. L., & Manheim, J. B. (2006). The one-step flow of communication. *The Annals of the American Academy of Political and Social Science*, 608(1), 213–232. <https://doi.org/10.1177/0002716206292266>
- Bimber, B., Flanagin, A. J., & Stohl, C. (2012). *Collective action in organizations: Interaction and engagement in an era of technological change*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511978777>

- Bond, S. (2024). *How AI deepfakes polluted elections in 2024*. NPR. <https://www.npr.org/2024/12/21/nx-s1-5220301/deepfakes-memes-artificial-intelligence-elections>
- Buccioli, A., & Zarri, L. (2013). *Lying in politics: Evidence from the US* (Working Paper Series, Department of Economics, University of Verona, No. 22). University of Verona. <https://linux2.dse.univr.it/home/workingpapers/wp2013n22.pdf>
- Burgoon, J. K., & Le Poire, B. A. (1993). Effects of communication expectancies, actual communication, and expectancy disconfirmation on evaluations of communicators and their communication behavior. *Human Communication Research*, 20(1), 67–96. <https://doi.org/10.1111/j.1468-2958.1993.tb00316.x>
- Burke, G., & Sunderman, A. (2024). *Brad Parscale helped Trump win in 2016 using Facebook ads. Now he's back, and an AI evangelist*. The Associated Press. <https://apnews.com/article/ai-trump-campaign-2024-election-brad-parscale-3ff2c8eba34b87754cc25e96aa257c9d>
- Carrasquillo, A. (2024). 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>
- Chatterjee, M. (2024). *What AI is doing to campaigns*. Politico. <https://www.politico.com/news/2024/08/15/what-ai-is-doing-to-campaigns-00174285>
- Chow, A. R. (2024). AI's underwhelming impact on the 2024 elections. *TIME*. <https://time.com/7131271/ai-2024-elections/>
- Cialdini, R. B., & Trost, M. R. (1998). Social influence: Social norms, conformity, and compliance. In D. T. Gilbert, S. T. Fiske, & G. Lindzey (Eds.), *The handbook of social psychology* (4th ed., pp. 151–192). McGraw-Hill.
- Coltin, J. (2024). How a fake, 10-second recording briefly upended New York politics. *Politico*. <https://www.politico.com/news/2024/01/31/artificial-intelligence-new-york-campaigns-00138784>
- Croson, R., Boles, T., & Murnighan, J. K. (2003). Cheap talk in bargaining experiments: Lying and threats in ultimatum games. *Journal of Economic Behavior & Organization*, 51(2), 143–159. [https://doi.org/10.1016/S0167-2681\(02\)00092-6](https://doi.org/10.1016/S0167-2681(02)00092-6)
- Dahl, D. W., Frankenberger, K. D., & Manchanda, R. V. (2003). Does it pay to shock? Reactions to shocking and nonshocking advertising content among university students. *Journal of Advertising Research*, 43(3), 268–280. <https://doi.org/10.1017/S0021849903030332>
- 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. (2024). *Data-driven campaigning and political parties: Five advanced democracies compared*. Oxford University Press. <https://doi.org/10.1093/oso/9780197570227.001.0001>
- Earl, J., & Kimport, K. (2011). *Digitally enabled social change: Activism in the internet age*. The MIT Press.
- Foos, F. (2024). The use of AI by election campaigns. *LSE Public Policy Review*, 3(3), 1–7. <https://doi.org/10.31389/lseppr.112>
- Fridkin, K. L., & Kenney, P. (2011). Variability in citizens' reactions to different types of negative campaigns. *American Journal of Political Science*, 55(2), 307–325. <https://doi.org/10.1111/j.1540-5907.2010.00494.x>
- Garimella, K., & Chauchard, S. (2024). How prevalent is AI misinformation? What our studies in India show so far. *Nature*, 630(8015), 32–34. <https://doi.org/10.1038/d41586-024-01588-2>
- Gracia, S. (2024). *Americans in both parties are concerned over the impact of AI on the 2024 presidential campaign*. Pew Research Center. <https://www.pewresearch.org/short-reads/2024/09/19/concern-over-the-impact-of-ai-on-2024-presidential-campaign/>
- Green, D. P., & Gerber, A. S. (2023). *Get out the vote: How to increase voter turnout* (5th ed.). Rowman & Littlefield. (Original work published 2004).
- Hackenburg, K., & Margetts, H. (2024). Evaluating the persuasive influence of political microtargeting with large language models. *PNAS: Proceedings of the National Academy of Sciences*, 121(24), e2403116121. <https://doi.org/10.1073/pnas.2403116121>
- Haselmayer, M. (2019). Negative campaigning and its consequences: A review and a look ahead. *French Politics*, 17(3), 355–372. <https://doi.org/10.1057/s41253-019-00084-8>

- Hersh, E. D. (2015). *Hacking the electorate: How campaigns perceive voters*. Cambridge University Press. <https://doi.org/10.1017/CBO9781316212783>
- Hindman, M. (2005). The real lessons of Howard Dean: Reflections on the first digital campaign. *Perspectives on Politics*, 3(1), 121–128. <https://doi.org/10.1017/S1537592705050115>
- Jay, M. (2010). *The virtues of mendacity: On lying in politics*. University of Virginia Press.
- Jost, J. T., Hennes, E. P., & Lavine, H. (2013). “Hot” political cognition: Its self-, group-, and system-serving purposes. In D. Carlston (Ed.), *The Oxford Handbook of social cognition* (pp. 851–875). Oxford University Press. <https://doi.org/10.1093/oxfordhb/9780199730018.013.0041>
- Jungherr, A. (2023). Artificial intelligence and democracy: A conceptual framework. *Social Media + Society*, 9(3), 1–14. <https://doi.org/10.1177/20563051231186353>
- 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>
- Jungherr, A., Schroeder, R., & Stier, S. (2019). Digital media and the surge of political outsiders: Explaining the success of political challengers in the United States, Germany, and China. *Social Media + Society*, 5(3), 1–12. <https://doi.org/10.1177/2056305119875439>
- Kahan, D. M. (2016). The politically motivated reasoning paradigm, part 1: What politically motivated reasoning is and how to measure it. In R. A. Scott & M. C. Buchmann (Eds.), *Emerging trends in the social and behavioral sciences* (pp. 1–16). John Wiley & Sons. <https://doi.org/10.1002/9781118900772.etrds0417>
- Kamal, R., & Kaur, J. (2025). Artificial intelligence and electoral decision-making: Analyzing voter perceptions of AI-driven political ads. In D. Goyal, B. Pratap, S. Gupta, S. Raj, R. R. Agrawal, & I. Kishor (Eds.), *Recent advances in sciences, engineering, information technology & management* (pp. 890–896). CRC Press.
- Karpf, D. (2012). *The moveon effect: The unexpected transformation of American political advocacy*. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780199898367.001.0001>
- Kreiss, D. (2011). Open source as practice and ideology: The origin of Howard Dean’s innovations in electoral politics. *Journal of Information Technology & Politics*, 8(3), 367–382. <https://doi.org/10.1080/19331681.2011.574595>
- Kreiss, D. (2012). *Taking our country back: The crafting of networked politics from Howard Dean to Barack Obama*. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780199782536.001.0001>
- Kreiss, D. (2016). *Prototype politics: Technology-intensive campaigning and the data of democracy*. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780199350247.001.0001>
- Kruschinski, S., Jost, P., Fecher, H., & Scherer, T. (2025). *Künstliche Intelligenz in politischen Kampagnen*. Otto Brenner Stiftung. <https://www.otto-brenner-stiftung.de/generative-ki-in-politischen-kampagnen/>
- Lau, R. R., & Rovner, B. (2009). Negative campaigning. *Annual Review of Political Science*, 12(1), 285–306. <https://doi.org/10.1146/annurev.polisci.10.071905.101448>
- Lin, Y.-C. (2023). 柯文哲推AI競選歌曲 談高虹安「凡事正其道而行」 [ko wen-je promotes AI campaign song and talks about KaoHung-an’s “doing everything in the right way”]. Radio Taiwan International. <https://www.rti.org.tw/news/view/id/2180407>
- Linville, D., & Warren, P. (2024). Digital yard signs: Analysis of an AI bot political influence campaign on X. Clemson University. https://open.clemson.edu/mfh_reports/7
- McCarthy, J. (2020). In U.S. Most oppose micro-targeting in online political ads. *Gallup Blog*. <https://news.gallup.com/opinion/gallup/286490/oppose-micro-targeting-online-political-ads.aspx>
- Muddiman, A. (2017). Personal and public levels of political incivility. *International Journal of Communication*, 11, 3182–3202.
- Murphy, G. L. (2002). *The big book of concepts*. The MIT Press.
- Nai, A. (2020). Going negative, worldwide: Towards a general understanding of determinants and targets of negative campaigning. *Government and Opposition*, 55(3), 430–455. <https://doi.org/10.1017/gov.2018.32>
- Nickerson, D. W., & Rogers, T. (2014). Political campaigns and big data. *The Journal of Economic Perspectives*, 28(2), 51–74. <https://doi.org/10.1257/jep.28.2.51>

- Nielsen, R. K. (2011). Mundane internet tools, mobilizing practices, and the coproduction of citizenship in political campaigns. *New Media & Society*, 13(5), 755–771. <https://doi.org/10.1177/1461444810380863>
- Ohtsubo, Y., Masuda, F., Watanabe, E., & Masuchi, A. (2010). Dishonesty invites costly third-party punishment. *Evolution and Human Behavior*, 31(4), 259–264. <https://doi.org/10.1016/j.evolhumbehav.2009.12.007>
- Peer, E., Rothschild, D., Gordon, A., Evernden, Z., & Damer, E. (2022). Data quality of platforms and panels for online behavioral research. *Behavior Research Methods*, 54(4), 1643–1662. <https://doi.org/10.3758/s13428-021-01694-3>
- Raj, S. (2024). How A.I. tools could change India's elections. *The New York Times*. <https://www.nytimes.com/2024/04/18/world/asia/india-election-ai.html>
- Raviv, S. (2025). When do citizens resist the use of AI algorithms in public policy? Theory and evidence. *The Journal of Politics*. <https://doi.org/10.1086/736362>
- Sanders, N. E., & Schneier, B. (2024). AI could still wreck the Presidential election. *The Atlantic*. <https://www.theatlantic.com/technology/archive/2024/09/ai-election-ads-regulation/680010/>
- Shirky, C. (2008). *Here comes everybody: The power of organizing without organizations*. The Penguin Press.
- Sifry, M. L. (2024). *How AI is transforming the way political campaigns work*. *The Nation*. <https://www.thenation.com/article/politics/how-ai-is-transforming-the-way-political-campaigns-work/>
- Simon, F. M., Altay, S., & Mercier, H. (2023). Misinformation reloaded? Fears about the impact of generative AI on misinformation are overblown. *Harvard Kennedy School Misinformation Review*, 4(5), 1–11. <https://doi.org/10.37016/mr-2020-127>
- Swenson, A. (2023). *Congressional candidate's voter outreach tool is latest AI experiment ahead of 2024 elections*. The Associated Press. <https://apnews.com/article/ai-chatbot-voters-election-2024-pennsylvania-db25cede35aa258d3e44563b517cc457>
- Swenson, A., Merica, D., & Burke, G. (2024). *AI experimentation is high risk, high reward for low-profile political campaigns*. The Associated Press. <https://apnews.com/article/artificial-intelligence-local-races-deepfakes-2024-1d5080a5c916d5ff10eadd1d81f43dfd>
- Ternovski, J., Kalla, J., & Aronow, P. M. (2022). The negative consequences of informing voters about deepfakes: Evidence from two survey experiments. *Journal of Online Trust & Safety*, 1(2), 1–16. <https://doi.org/10.54501/jots.v1i2.28>
- Tomić, Z., Damnjanović, T., & Tomić, I. (2023). Artificial intelligence in political campaigns. *South Eastern European Journal of Communication*, 5(2), 17–28. <https://doi.org/10.47960/2712-0457.2.5.17>
- Turrow, J., Carpini, M. X. D., Draper, N., & Howard-Williams, R. (2012). *Americans roundly reject tailored political advertising: At a time when political campaigns are embracing it*. Annenberg School for Communication, University of Pennsylvania. <https://repository.upenn.edu/handle/20.500.14332/2053>
- Tyler, J. M., Feldman, R. S., & Reichert, A. (2006). The price of deceptive behavior: Disliking and lying to people who lie to us. *Journal of Experimental Social Psychology*, 42(1), 69–77. <https://doi.org/10.1016/j.jesp.2005.02.003>
- van Kleef, G. A., Wanders, F., Stamkou, E., & Homan, A. C. (2015). The social dynamics of breaking the rules: Antecedents and consequences of norm-violating behavior. *Current Opinion in Psychology*, 6, 25–31. <https://doi.org/10.1016/j.copsyc.2015.03.013>
- Verma, P., & Vynck, G. D. (2024). AI is destabilizing 'the concept of truth itself' in 2024 election. *Washington Post*. <https://www.washingtonpost.com/technology/2024/01/22/ai-deepfake-elections-politicians/>
- Weiss-Blatt, N. (2021). *The techlash and tech crisis communication*. Emerald Publishing.
- West, D. M., & Allen, J. R. (2020). *Turning point: Policymaking in the era of artificial intelligence*. Brookings Institution Press.
- Williams, D. (2023). The case for partisan motivated reasoning. *Synthese*, 202(89), 1–27. <https://doi.org/10.1007/s11229-023-04223-1>
- Zhang, B., & Dafoe, A. (2019). *Artificial intelligence: American attitudes and trends*. Center for the Governance of AI University of Oxford. <https://doi.org/10.2139/ssrn.3312874>