



The Vigilant Public: Awareness, Skepticism, and Perceived Democratic Impact of AI-Generated Deepfakes Among Indian Voters

Shubham Bhatia¹ Prof. Vikram Kaushik²

Department of Mass Communication, GJUS&T, Hisar, Haryana, India, PIN: 125001

Department of Mass Communication, GJUS&T, Hisar, Haryana, India, PIN: 125001

* Corresponding author: Bhatiasubham718@gmail.com, Mobile: 9805627428 • ORCID: 0009-0000-4612-6264

Abstract

The rapid diffusion of generative artificial intelligence (AI) has made synthetic political media deepfakes and cheapfakes an increasingly salient concern for democratic communication. This study examines public awareness, perception, skepticism, and self-reported behavioral response to AI-generated political content among a diverse sample of 404 respondents. Using a structured survey instrument covering demographic profile, awareness, individual perception, decision-making impact, media skepticism, and democratic participation, the study finds near-universal awareness of AI (99.5%) but comparatively lower recognition of specific terminology such as 'deepfake' and 'cheapfake' (85.6%). A recurring third-person pattern emerges: respondents report personal resilience to influence (61.9% say their own opinions are unaffected) while simultaneously reporting that AI-generated media erodes trust and discourages public dialogue at a societal level (69.8–73.0% agreement). The findings suggest a public that is highly aware but deeply skeptical, one that is adapting through heightened vigilance rather than blind acceptance, with clear implications for platform labeling policy, media literacy interventions, and electoral regulation.

Keywords: Deepfakes, cheap fakes, political misinformation, artificial intelligence, voter perception, democratic participation

1. Introduction

Generative artificial intelligence has transformed the production of AI-generated synthetic media from a specialized technical capability of individuals into a widely accessible consumer tool. Photorealistic images, cloned voices, and manipulated videos like 'deepfakes' and 'cheapfakes' can now be created within seconds using freely available applications on the internet. While such tools have legitimate uses in entertainment, education, and creative industries, their use in political scenarios and political processes like elections has raised urgent questions about voter trust, democratic values, and the integrity of democratic discourse. Recent election cycles all over the world have provided concrete illustrations of this risk. Examples like the AI-generated robocall incident involving former U.S. president Joe Biden were used to discourage primary voters from turning out; political figures in countries like India, the UK, Türkiye, and the Philippines have been targeted with these AI-generated fabricated audio and video. At the same time, empirical audits of election-related deepfake studies suggest that the overall volume of this synthetic content remains small compared to non-AI 'cheap' misinformation. The recent spread of this malicious AI-generated content has documented enhanced harm to public trust and made people vulnerable to this misinformation. This study of voter



awareness and perception especially focuses on how citizens perceive this AI-generated content, how it impacts their trust, and how they respond to this misinformation.

The present study investigates a sample of largely young, educated, and rural and urban respondents in India. As the reach of the internet has spread throughout the country, social media has changed political campaigning in India through various techniques. In this survey study, eligible voters were used as the sample, as they understand and react to AI-generated political content. This study specifically addresses five interlinked research questions: (1) How aware are respondents of AI and its capacity to generate synthetic content? (2) What are respondents' individual perceptions of AI-generated political content in terms of engagement? (3) How does exposure to such content affect individual decision-making and opinion formation? (4) How skeptical are respondents of AI-generated political media? (5) Does the presence of AI-generated media affect respondents' willingness to participate in democratic processes? Rather than simply relying on descriptive percentages, this study tests response patterns across 404 respondents to determine whether the reported attitudes are systematic and statistically robust, and to uncover any gap between respondents' sense of personal resilience and their broader concern for democratic discourse, a pattern well documented in communication research as the third-person effect.

2. Literature Review

2.1. The Prevalence and Political Use of Deepfakes.

Research tracking real-world deep-fake incidents has found a diversification of political uses [1]. In their study, they found that the incidents from mid-2023 to mid-2024 identify electioneering, character assassination, and non-consensual targeting of women in politics as dominant categories, across dozens of countries. Walker et al. [2] dedicated incident databases, such as the Political Deepfakes Incidents Database, demonstrate a growing consensus among academics that systematic tracking is needed to move the discussion beyond anecdote. Kishwar et al. [3] suggest that the 2024 electoral cycle was much affected by the non-AI 'cheap fake' misinformation, which was used much more frequently than artificial intelligence-generated content.

2.2. Effects on Deception, Trust, and Political Attitudes

Experimental evidence on the effect of deepfakes on audiences is more nuanced than public alarm might suggest. In a large representative-sample experiment in the United Kingdom, Vaccari and Chadwick [4] found that respondents who were exposed to deepfake videos were more likely to be uncertain about what they had seen than to be straightforwardly deceived by it, and that it was this heightened uncertainty rather than direct deception that reduced viewers' trust in news encountered on social media. This is a process of generalized indeterminacy and cynicism that deepfakes add to civic culture, they write, with the central risk of political deepfakes moving away from mass persuasion and toward a more corrosive erosion of general trust in information. Dobber et al. [5] complement this by producing an original political deepfake and testing its effects in an online experiment (N = 278). In their study they found that the attitudes toward the depicted politician were significantly lower after exposure, but attitudes toward the politician's party were largely unaffected; critically, they also find that microtargeting the deepfake at specific, susceptible audience segments amplified its persuasive impact relative to untargeted exposure, suggesting that the political risk of deepfakes may be concentrated within particular subgroups of the electorate rather than distributed evenly across it.

2.3. The Third-Person Effect in Misinformation and Deepfake Perception

A well-known phenomenon in communication research is the tendency for people to assume that persuasive or harmful media content affects others more than themselves, which is referred to as the third-person effect [6] and extensively studied in the setting of misinformation. Li et al. [7] conducted a meta-analysis of 30 studies and concluded that exposure to distorted information induces a strong third-person perception across contexts ($r = .634$). The perceptual gap was significantly linked to behavioral outcomes. It positively predicted corrective action, but it was negatively associated with the intention to share questionable content. Using survey data from 1,111 Chinese citizens, Tang et al. [8] similarly found that greater exposure to fake news and higher perceived information overload were both associated with stronger third-person perceptions, and notably that respondents who believed fake news affected others more than themselves were less likely to support stricter regulatory controls on fake news. This latter finding bears directly on the present study's observation that respondents reporting high personal resilience to AI-generated political content simultaneously report considerable concern about its effects on democratic discourse at large.

2.4. Third-Person Perceptions, Corrective Action, and Regulatory Support

Li and Yan [9] further shed light on the behavioral consequences of this self-other perceptual gap, by employing structural equation modeling on survey data from 1,500 Chinese medical students to demonstrate that stronger third-person perceptions of digital misinformation were negatively associated with respondents' willingness to engage in corrective actions, such as rebutting false claims or promoting accurate information, and that both self-efficacy and collectivist orientation positively predicted corrective engagement. Consistent with the meta-analytic evidence in [7], this body of work indicates that the third-person gap observed in the present study's findings is not purely descriptive: a larger perceived gap between self and society may plausibly suppress individual motivation toward corrective or vigilant behavior, even as collective concern about misinformation is high.

2.5. Regulatory Responses to Political Deepfakes

Policy responses to political deepfakes in democratic countries remain fragmented. Chawki [10], in his study, reviewed the framework of the American federal and state-level legal responses to AI-generated content like deepfake videos, images, and voice cloning and concluded that, despite a growing number of laws or solutions like labeling of AI-generated political content, the regulatory and detection capacities continue to lag behind. In the Indian context, regulatory safeguards have visibly lagged behind political adoption of the technology: Sahoo [11] documents many major incidents involving Indian national parties in which the use of deepfakes as an official campaign tool during political campaigns goes viral on various social media platforms. The use of such technologies in that way emerged largely because of the absence of adequate state regulation. Because of this unsettled regulatory landscape, this study tries to explore awareness, perception, and what Lehman thinks about the use of AI-generated content in political scenarios.

3. Methodology

3.1. Research Design

This study uses the quantitative, cross-sectional survey design method to collect data from eligible respondents to capture descriptive and inferential insights into their awareness, perceptions, skepticism, and self-reported behavioral responses to AI-generated political content. The research design is exploratory-descriptive in nature and does not attempt to establish causal relationships between variables. A total of (N=404) valid responses were collected. Respondents were drawn from a mixed population of college and

university students, research scholars, educators, and individuals in other occupational categories. The sample contained 18–29-year-olds (52.5%) and those with a graduate degree or higher (89.9%) to reflect the relative access of those populations to the survey instrument and its digital reach. With this composition, the sample should be regarded as broadly representative of the digitally active and educated sections of the population rather than the general electorate.

3.2. Instrument

The survey instrument consisted of six sections: (i) demographic information (gender, age, qualification, occupation, location); (ii) awareness of AI and AI-generated synthetic media, measured through six dichotomous (yes/no) items; (iii) individual perceptions of AI-generated political content, measured on six five-point Likert-type items; (iv) impact of AI-generated media on individual decision-making, measured through four Likert items; (v) skepticism and trust toward AI-generated political content, measured through six Likert items; and (vi) impact on democratic participation, measured through three Likert items. All Likert items used a five-point agreement scale (Strongly Agree, Agree, Neutral, Disagree, Strongly Disagree).

3.3. Data Analysis

- The chi-square goodness-of-fit tests were used on the six awareness items to test whether the observed aware/not-aware split differed significantly from a 50/50 chance distribution.
- The Chi-square goodness-of-fit tests were used on all nineteen Likert-type items to test whether the observed five-category response distribution differed significantly from a uniform distribution (20% per category), indicating that responses reflect genuine, non-random attitudes rather than indifference.
- Two-proportion z-tests were used to compare the proportion of respondents who agreed (Strongly Agree + Agree) against the proportion who disagreed (Disagree + Strongly Disagree) for each Likert item, testing the null hypothesis that agreement and disagreement are equally likely (50/50) among respondents who took a position.
- Weighted mean agreement scores (1 = Strongly Disagree to 5 = Strongly Agree) were computed for each Likert item to allow ranked comparison of attitude strength across all nineteen items.

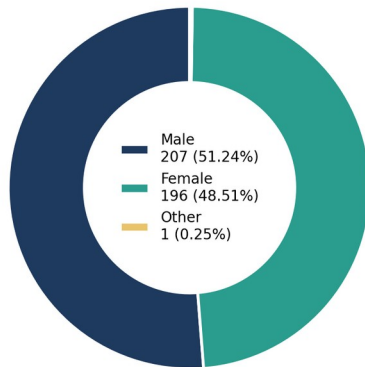
All tests were conducted at a significance threshold of $\alpha = .05$ using Python (SciPy 1.17).

4. Results and Analysis

4.1. Demographic Profile of Respondents

The sample used in this study consisted of 404 respondents with near gender parity (51.2% male, 48.5% female, 0.25% other). The respondent population included younger respondents, with over half (52.5%) aged 18–29 and was predominantly urban (63.1%) and highly educated, with 89.9% holding or pursuing at least a graduate degree. In the total sample, the students formed the largest occupational group (44.3%), followed by respondents in other occupations (21.5%), research scholars (19.8%), and educators (14.4%).

Gender Distribution of Respondents (N=404)



Rural-Urban Distribution (N=404)

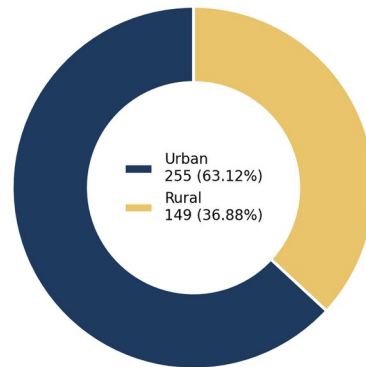


Figure 1. Gender and rural-urban distribution of respondents (N = 404).

Age Distribution of Respondents (N=404)

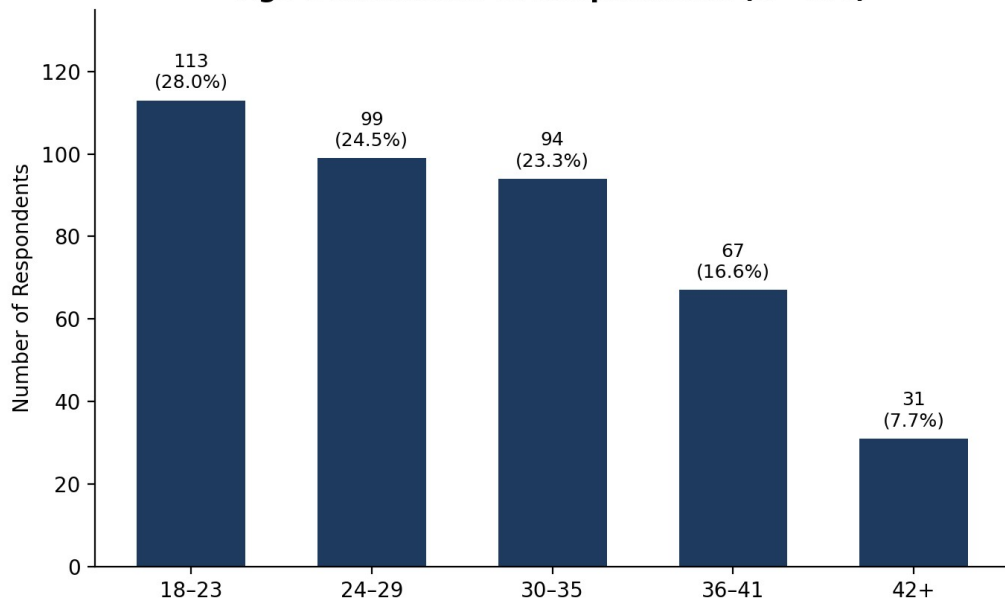


Figure 2. Age distribution of respondents (N = 404).

Sample Composition: Qualification and Occupation (N=404)

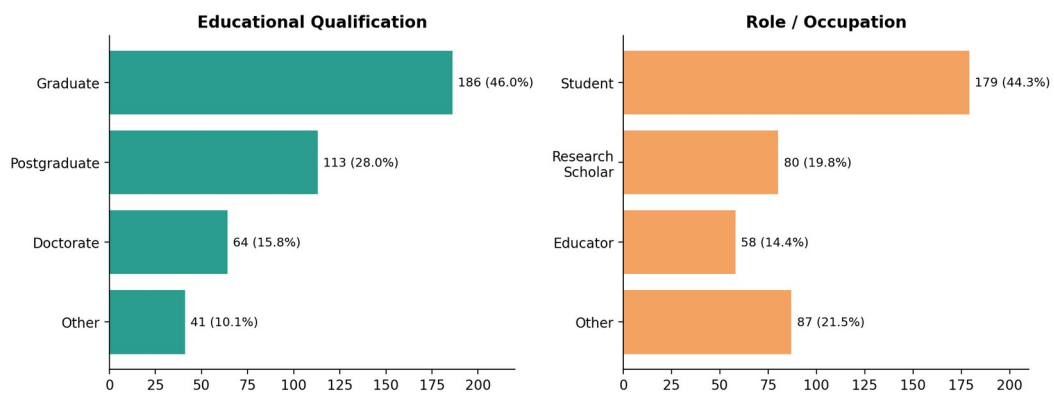


Figure 3. Educational qualification and occupational role of respondents (N = 404).

4.2. Awareness of AI-Generated Synthetic Media

Awareness of artificial intelligence was near-universal: 99.5%. The maximum number of respondents reported knowing what AI means and had heard of its use in everyday technology. They also believed AI can generate videos, images, or text content quickly. Awareness of specific terms like ‘deepfake’ and ‘cheap fake’ in each qualifications group was comparatively lower at 85.6%, and awareness that AI can be used to generate cloned voices stood at 92.8%. Awareness among respondents about the use of AI-generated content in political campaigns was high at 97.5%. Chi-square tests show that all awareness items are significantly different from a chance (50/50) distribution (range: 205.31–396.04, all $p < .001$), suggesting that awareness in this sample is a real, robust trait and not the result of random patterns of response.

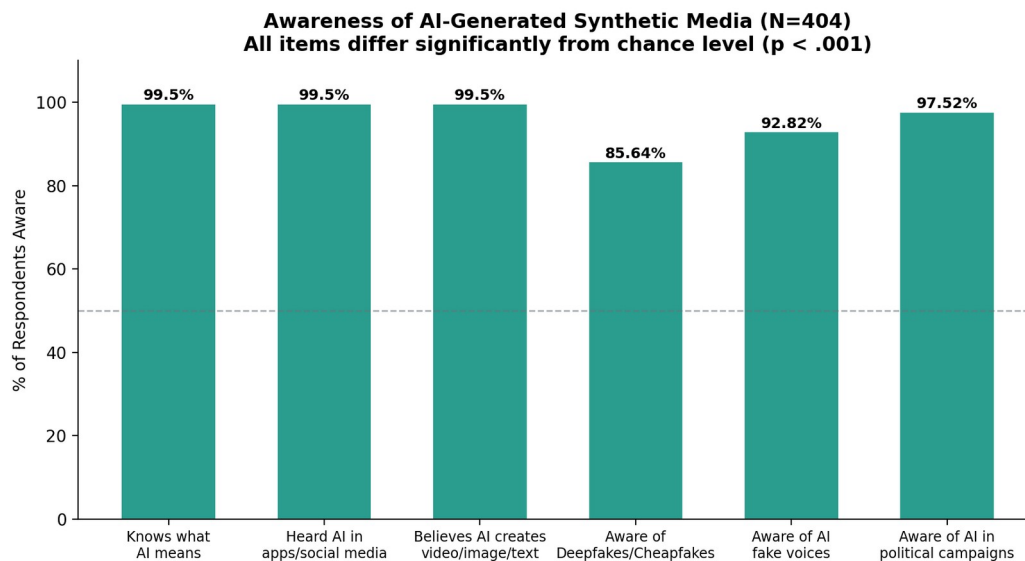


Figure 4. Percentage awareness across six AI/synthetic-media awareness items (N = 404).

4.3. Individual Perceptions of AI-Generated Political Content

According to the findings, the largest number of respondents reported having encountered AI-altered content on online social media platforms (91.3% agreement, mean = 4.48/5). They also expressed their concern about misinformation, consent, and manipulation associated with AI-generated media (92.6% agreement, mean = 4.50/5, $z = 19.34$, $p < .001$). At the same time, three-quarters of respondents (75.0%) believe that AI-generated content is engaging and creative, suggesting that these types of media-related concerns and appreciation coexist rather than being mutually exclusive. The majority of respondents (58.4%) rejected the statement that AI-generated media has no major impact on the trustworthiness of political information (mean = 2.73, $z = -5.78$, $p < .001$). On the other hand, 67.3% agreed that AI media is used more to spread misinformation and manipulate than to educate the public in political participation. Another section of respondents (69.6%) believes that trusting AI-generated content more and having less harm on their critical thinking when it is clearly labeled by their creators and platforms where it is published ($z = 13.84$, $p < .001$). This specific insight highlights that labeling and transparency-like policies related to such AI-generated content act as a potentially effective trust-restoration mechanism among the public.

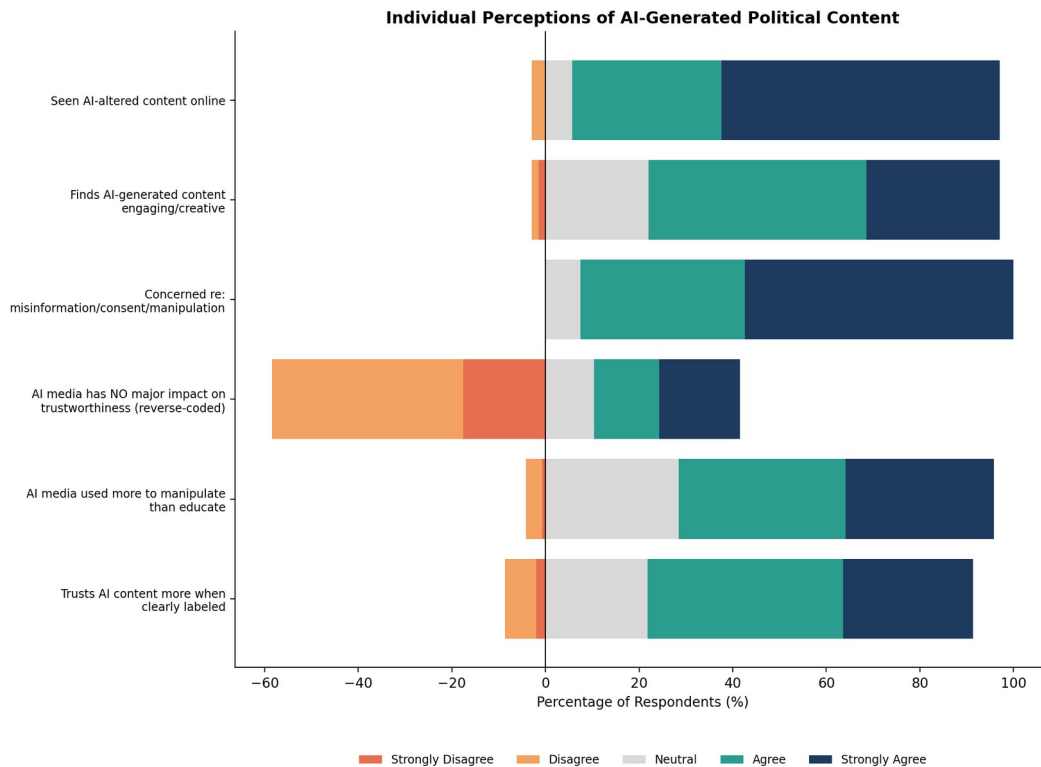


Figure 5. Distribution of responses on individual perception items (N = 404).

4.4. Impact on Individual Decision-Making

The Findings on decision-making reveal a complex, at times internally conflicted, pattern. While 48.8% of respondents agreed that content produced by AI does not make it harder to verify truth online. The difference in the dataset from the 50/50 baseline did not reach significance in the expected direction as strongly as other items, $z = 4.45$, $p < .001$, although the agree/disagree gap here was the narrowest of all nineteen items. Most respondents (56.4%) simultaneously admitted that AI-generated political videos and images influence their opinions even when they suspect the content might be fake ($z = 7.60$, $p < .001$). On the other hand, a larger majority (61.9%) believes that AI-generated content does not change their opinions on political issues or candidates ($z = 9.78$, $p < .001$), suggesting many respondents distinguish between transient influence and durable opinion change. Strikingly, 70.8% of respondents agreed that the presence of AI-generated media has made them more cautious in their decision-making overall (mean = 3.91, $z = 14.40$, $p < .001$), with very low outright disagreement (7.4%).

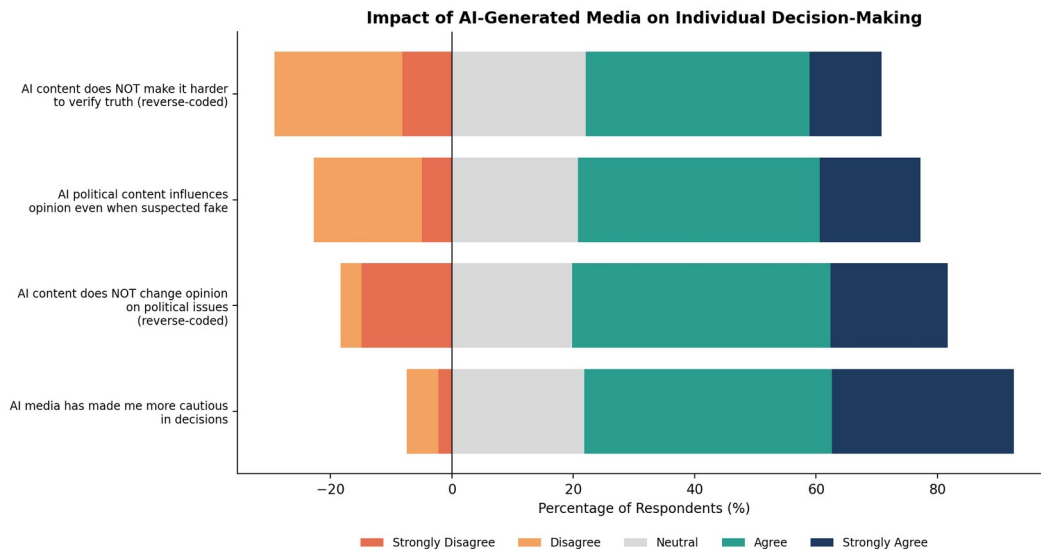


Figure 6. Distribution of responses on the impact of AI-generated media on individual decision-making (N = 404).

4.5. Skepticism and Trust Toward AI-Generated Political Content

As most of the respondents are educated and aware of the presence of AI-generated content on online platforms, skepticism emerged as the dominant orientation across this section. Nearly 80% of respondents (79.7%) reported being skeptical of AI-generated media content. They believe that it is important to verify such content through trusted sources (mean = 4.10, $z = 17.12$, $p < .001$). From these respondents the 71.5% reported have doubt the reliability of AI-generated political media even it is clearly labelled as AI-made ($z = 13.65$, $p < .001$) a finding that qualifies the labelling-trust result of Section 4.3, suggesting that labeling increases relative but not absolute trust. Most (67.8%) agreed that content produced by AI does more harm than good for democratic society ($z = 12.44$, $p < .001$), and more than half (56.9%) said content produced by AI has led them to question the authenticity of information they use personally or professionally. In line with this general skepticism, a majority (55.9%) disagreed that existing laws and policies are sufficient to prevent the misuse of AI-generated synthetic media, compared to 31.4% who agreed ($z = -5.27$, $p < .001$) suggesting a widely perceived regulatory gap.

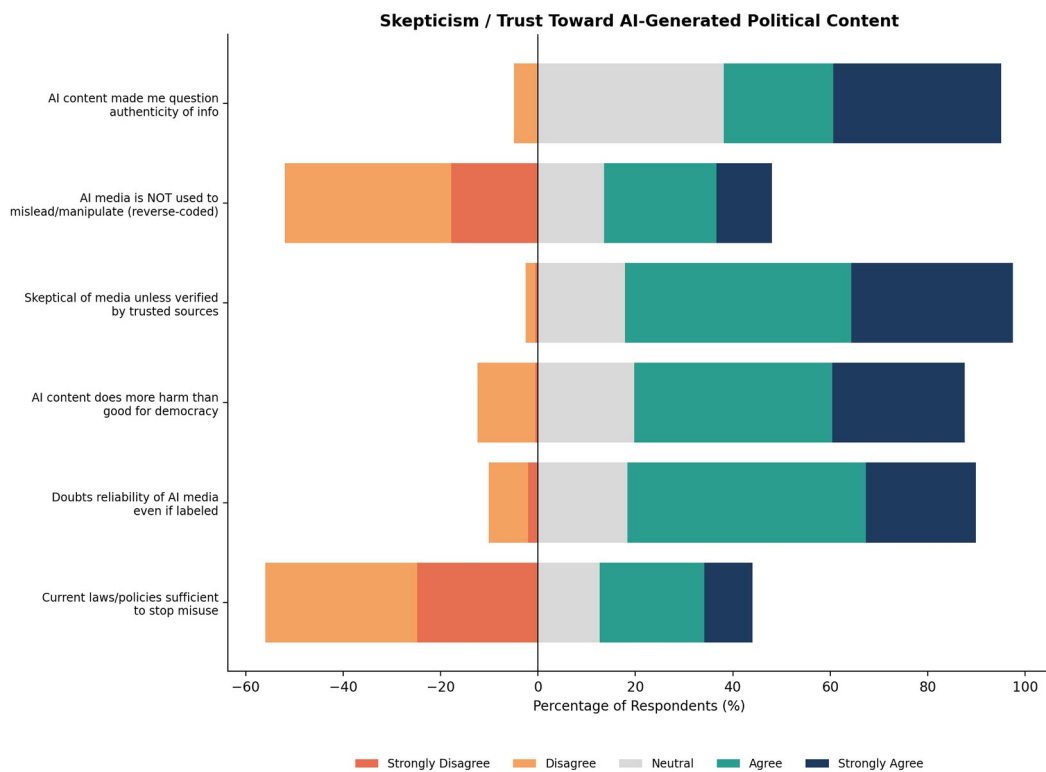


Figure 7. Distribution of responses on skepticism and trust toward AI-generated political content (N = 404).

4.6. Impact on Democratic Participation

A revealing disconnect appears between respondents' personal and societal-level assessments. A significant number of respondents (67.6%) have confidence that AI-generated media does not affect their willingness to participate in democratic processes like voting and political discussion ($z = 11.17$, $p < .001$). In comparison, the larger majority (73.0%) agreed that exposure to such AI-generated manipulated media makes it difficult to trust political information. Because of that they did not engage confidently in democratic decision-making ($z = 14.01$, $p < .001$). 69.8% of respondents also agreed that the spread of AI-generated manipulated media impacts negatively and discourages meaningful public dialogue. This result weakens their overall belief in democratic participation ($z = 12.88$, $p < .001$). This pattern of self-reported personal resilience alongside pronounced concern about collective, societal-level erosion of trust is consistent with the third-person effect.

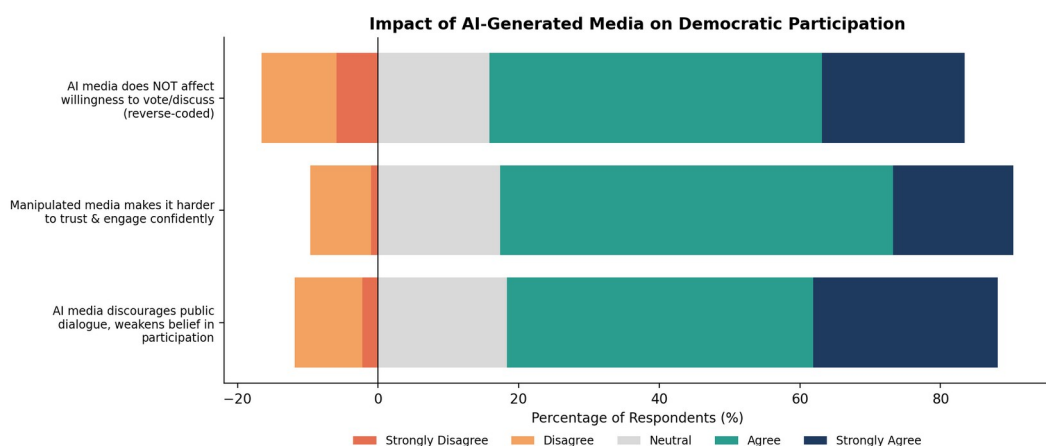


Figure 8. Distribution of responses on the perceived impact of AI-generated media on democratic participation (N = 404).

4.7. Comparative Summary Across Attitude Items

Figure 9 ranks all 19 Likert-type items by mean agreement score. The highest scoring items are related to worrying about misinformation (mean = 4.50) and having seen AI-manipulated content (mean = 4.48), while the lowest scoring items are reverse-worded statements expressing confidence in current laws being enough (mean = 2.61) and that AI-generated media has no big effect on trustworthiness (mean = 2.73). The ranking reveals that respondents' strongest and most statistically robust beliefs concern risks and regulatory shortcomings of AI-generated media. Conversely, agreement is comparatively muted – but non-random – regarding items that suggest personal invulnerability.

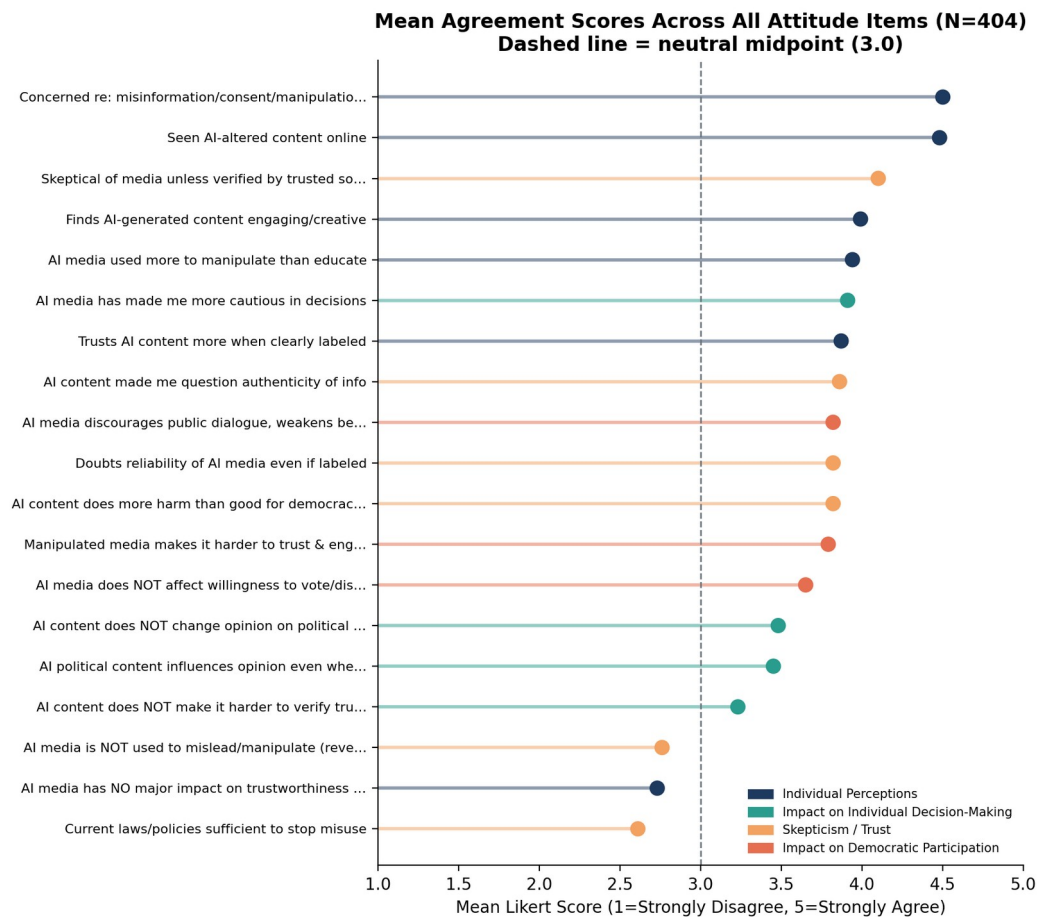


Figure 9. Mean agreement scores (1–5 scale) across all nineteen attitude items, ranked from lowest to highest

Summary of Inferential Test Results

Table 1 summarizes the chi-square goodness-of-fit and two-proportion z-test results for the items with the strongest and most policy-relevant findings across each section of the survey.

Item (Section)	% Agree	% Disagree	χ^2 (GOF) / p	z (Agree vs Disagree) / p
Seen AI-altered content online (Perceptions)	91.3	3.0	523.2 / p<.001	18.29 / p<.001
Concerned re: misinformation/manipulation (Perceptions)	92.6	0.0	522.8 / p<.001	19.34 / p<.001
AI media has real impact on trustworthiness (Perceptions, reverse)	31.2	58.4	116.6 / p<.001	-5.78 / p<.001
More cautious in decisions (Decision-making)	70.8	7.4	216.4 / p<.001	14.40 / p<.001
Skeptical unless verified by trusted sources (Skepticism)	79.7	2.5	320.7 / p<.001	17.12 / p<.001

Item (Section)	% Agree	% Disagree	χ^2 (GOF) / p	z (Agree vs Disagree) / p
Current laws sufficient to stop misuse (Skepticism, reverse)	31.4	55.9	61.9 / p<.001	-5.27 / p<.001
Manipulated media reduces trust/confidence to engage (Democratic)	73.0	9.7	363.1 / p<.001	14.01 / p<.001
Media discourages dialogue, weakens participation belief (Democratic)	69.8	11.9	206.0 / p<.001	12.88 / p<.001

Table 1. Selected chi-square goodness-of-fit and two-proportion z-test results across survey sections (N = 404; all reported p-values < .001).

5. Discussion

The respondent base from the sample is highly informed about the term artificial intelligence in general, moderately familiar with specific synthetic-media terminology like deepfakes and cheap fakes. They are persistently and statistically suspicious of AI-produced political content. Virtually all of the attitude questions in the questionnaire are statistically tested and differ significantly from the random distribution, suggesting that the respondents have well-defined opinions and are not indifferent. This is in line with the general claim made by Jacobsen and Simpson [1], according to which deepfakes have already become fully integrated into modern politics and media worries, so people's reactions to them are guided by existing concerns about image manipulation rather than misunderstanding of a new technology.

There are three conclusions in particular that appear to be very significant. First, the existence of both high engagement of the respondents (75% consider AI-produced political content engaging and creative) and high concern about misinformation (92.6%). This means that an understanding of the creative potential of this content does not lead to indifference to the risks of using it. These two attitudes are apparently independent. Secondly, the importance respondents attach to labeling, 69.6% of whom trust labeled content, but still, 71.5% of them doubt its reliability is completely consistent with the results of experimental studies of AIGC disclosure, which has found that labeling AI-generated material produces, at best, a modest and inconsistent effect on perceived credibility rather than a reliable boost in trust [12]. This means that transparency measures such as labeling content produced by AI may shift relative trust levels, but not the basic suspicion uncovered throughout this survey. Third, and importantly, the discrepancy between personal and societal-level assessments (67.6% report no personal effect on their willingness to engage while 69.8-73.0% report broader erosion of dialog and trust) closely resembles the third-person effect well-documented in previous misinformation research, where people consistently rate themselves as less susceptible to media influence than they rate others or society at large [6], [8]. This third-person pattern has a double-edged implication. On the other hand, self-perceived resilience may serve as a proxy for critical-thinking capacity in an educated, digitally literate sample. Conversely, previous third-person effect research warns that a salient self-other gap can diminish individuals' motivation to engage in corrective behavior: structural equation modeling on a large sample of Chinese medical students found that stronger third-person perceptions of digital misinformation were negatively associated with respondents' willingness to rebut false claims or promote accurate information [9].

However, it seems as if there is a contradiction in the results of the decision-making part of the research, which states that even while admitting the fact that the AI-produced political information influences their opinion, and they assume that it might be faked, 56.4% of the respondents still say that they do not let such

information affect their opinion on political matters or politicians. Such contradictory results have been reported before. Vaccari and Chadwick [4], in their experimental study, found that a representative sample of the UK population, deepfakes create doubts rather than deceive people, and those doubts decrease trust in the news posted on social networks. Lewandowsky et al. [13] also found that misinformation continues to influence people's memory even after it has been corrected, called the continued influence effect. Similar dynamics might explain the results of the current study. In addition, Dobber et al. [5] found in experiments that a political deepfake might lower attitudes towards the politician without having any effect on attitudes towards his or her party, and that micro-targeting certain groups of respondents exacerbates this effect.

The large proportion of the population in the survey believes that the current legislation is ineffective and inefficient in addressing the problem of use of AI-generated synthetic media (55.9% disagreement vs 31.4% agreement). The findings fit well within the context of the existing regulation literature, which also identifies the current legislative response to political deepfakes as fragmented and reactive [10]. The situation in India specifically is also consistent with this view, as reported by Sahoo [11]. In his study, he finds that political deepfakes were used as a campaign technique in India by a major political party without any regulation and only with the government warning after the fact to label and quickly remove any election-related synthetic content. In combination with the fact that content labeling enjoys some degree of trust among respondents, this leads one to assume that the two measures might be seen as complementary, not substitutable solutions.

Figure 10 presents all the above findings within a single conceptual framework and outlines the general process from awareness to concern, skepticism, and cautious behavior considering the two moderators.

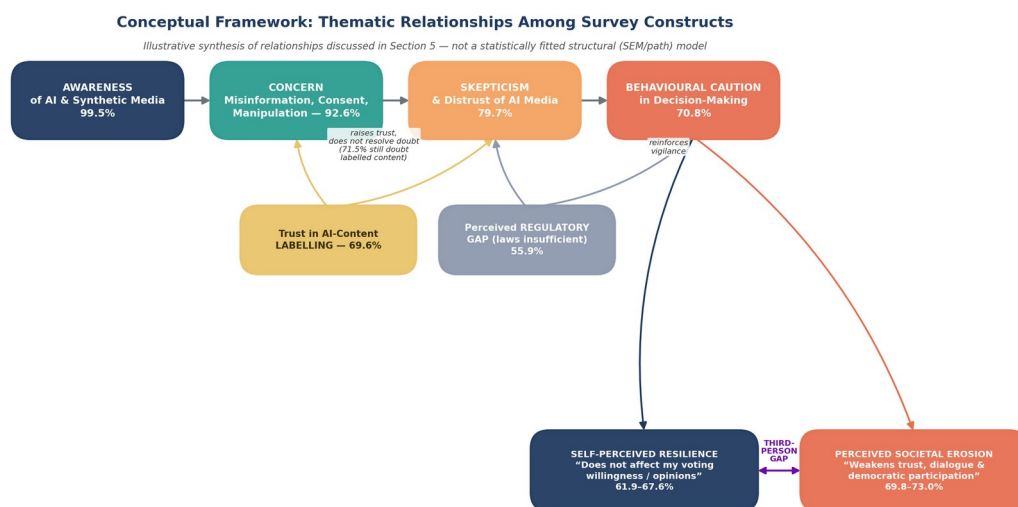


Figure 10. Conceptual framework synthesizing thematic relationships among the study's key constructs (N = 404).

6. Conclusion

This study examines awareness and perception of AI-generated political deepfakes among a sample of 404 respondents. In order to determine whether the attitudinal patterns presented are statistically robust or mere descriptions. The latter seems to be the case because awareness of AI is very high, and every single attitudinal statement showed a significant difference from random, proving the existence of definite opinions on the matter. This result is particularly valuable in light of the recent political deepfake scandal in India. The demographics of the participants (young, urban, and highly educated) are consistent with the group most

exposed to fake political material in the context of India's 2024 general elections, the biggest democratic exercise in the world, where over 75% of Indians were exposed to political deepfakes and one quarter took them to be true [14], [15]. WhatsApp, which is used by more than 400 million Indians and which was also labeled as the main vehicle of election-related fake news, was central to this exposé [16]; an experience that aligns with the current study's finding that 91.3% of participants have been exposed to AI-tweaked information. In this context, the high degree of skepticism (79.7%) expressed in the form of not trusting anything unless proven true and the high number of people who believe that AI media can do more harm than good for democracy (67.8%) appear to be less about abstract fear and more about an actual experience. The fact that more than half (55.9%) of respondents feel current laws and policies do not adequately address issues related to curbing abuse of AI-generated synthetically manipulated media also proves true in the case of the Indian regulatory history. According to [11], political deepfakes had already been used by a prominent national party in its official multilingual campaign effort even prior to the existence of any kind of regulation for the same, and the amendment to the Indian Information Technology Rules defining "synthetically generated information" took place only in November 2025, i.e., post the 2024 election cycle when these perceptions were formed.

The third-person gap observed in this study between personal resilience (61.9%-67.6%) and high levels of anxiety about the erosion of trust and dialogue on a societal level (69.8%-73.0%) also relates to an Indian communication system that is specific to political content. In India, political discourse is largely spread through private, secure WhatsApp groups of family and community, not on open media channels, and this may give people a feeling of security on a personal level, despite understanding the damaging effects of manipulation in the overall information ecosystem [11]. Together with the fact that more than half of India's internet users reside in rural areas (whereas the sample used for this survey was predominantly urban – only 36.9% rural), this means that the level of awareness and skepticism reported in this survey is most likely a ceiling, not the national average.

Far from inducing uncritical acceptance or panic, the growth of AI-driven political content seems, in this case, to be breeding a culture of epistemic vigilance. In the case of India especially, where elections are waged on an unprecedented scale, in more than one language, and through AI-driven content [11], ensuring that this vigilance is sustained and channeled constructively in the coming years will be contingent on three simultaneous efforts: Enforcement by Election Commission and Ministry of Electronics and Information Technology matching the pace of the campaigns and not falling behind them; the consistent enforcement of the labeling and traceability requirements of the amended IT Rules on the vernacular, WhatsApp-first platforms where Indians actually consume political content; and media literacy campaigns that aim to address the self-other perception problem, and not just knowledge, particularly among rural and older voters who do not belong to the digital-savvy cohort sampled in this study.

Acknowledgements

The authors wish to thank Prof. Vikram Kaushik for his valuable insights in this project. The author also sincerely appreciates Dr. Yukti Dhadwal for their generous help during the various stages of this study

Funding

The author(s) received no financial support for the research, authorship, and/or publication of this article.

Conflict of Interest

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Data Availability Statement

The dataset analyzed in this study is not publicly available, as it is currently being utilized for ongoing doctoral research. To protect the integrity and confidentiality of the PhD thesis work, direct access to the raw data cannot be provided at this time.

AI Usage Disclosure

No AI tool was used to generate the content.

References

- [1] B. N. Jacobsen and J. Simpson, "The tensions of deepfakes," *Information, Communication & Society*, vol. 27, no. 6, pp. 1095–1109, 2023, doi: 10.1080/1369118X.2023.2234980.
- [2] C. P. Walker, D. S. Schiff, and K. J. Schiff, "Merging AI incidents research with political misinformation research: Introducing the Political Deepfakes Incidents Database," *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, no. 21, pp. 23053–23058, 2024, doi: 10.1609/aaai.v38i21.30349.
- [3] S. D. Kishwar, A. Tripathi, S. Khatoon, D. Poddar, and B. Khurana, "Regulating deep fakes and synthetic media: Privacy, policy and global regulatory challenges," *Journal of Data Protection & Privacy*, vol. 8, no. 1, pp. 78–96, 2025.
- [4] C. Vaccari and A. Chadwick, "Deepfakes and disinformation: Exploring the impact of synthetic political video on deception, uncertainty, and trust in news," *Social Media + Society*, vol. 6, no. 1, pp. 1–13, 2020, doi: 10.1177/2056305120903408.
- [5] T. Dobber, N. Metoui, D. Trilling, N. Helberger, and C. de Vreese, "Do (microtargeted) deepfakes have real effects on political attitudes?," *The International Journal of Press/Politics*, vol. 26, no. 1, pp. 69–91, 2021, doi: 10.1177/1940161220944364.
- [6] W. P. Davison, "The third-person effect in communication," *Public Opinion Quarterly*, vol. 47, no. 1, pp. 1–15, 1983, doi: 10.1086/268763.
- [7] X. Li, W. Lyu, S. M. Salleh, et al., "Combating distorted information through third person perception: Advances in a meta-analysis," *Current Psychology*, vol. 45, Art. no. 673, 2026, doi: 10.1007/s12144-026-09212-4.
- [8] S. Tang, L. Willnat, and H. Zhang, "Fake news, information overload, and the third-person effect in China," *Global Media and China*, vol. 6, no. 4, pp. 492–507, 2021, doi: 10.1177/20594364211047369.
- [9] Z. Li and J. Yan, "Does a perceptual gap lead to actions against digital misinformation? A third-person effect study among medical students," *BMC Public Health*, vol. 24, no. 1, pp. 1–13, 2024, doi: 10.1186/s12889-024-18763-9.
- [10] M. Chawki, "Navigating legal challenges of deepfakes in the American context: A call to action," *Cogent Engineering*, vol. 11, no. 1, Art. no. 2320971, 2024, doi: 10.1080/23311916.2024.2320971.
- [11] S. Sahoo, "A critical political economy of campaign deepfakes in Indian elections," in *Critical Political Economy and Southern Approaches to AI*, P. Raghunath, Ed. Cham, Switzerland: Palgrave Macmillan, 2026, doi: 10.1007/978-3-032-14090-6_4.
- [12] F. Li, Y. Yang, and G. Yu, "Nudging perceived credibility: The impact of AIGC labeling on user distinction of AI-generated content," *Emerging Media*, vol. 3, no. 2, pp. 275–304, 2025, doi: 10.1177/27523543251317572.
- [13] S. Lewandowsky, U. K. H. Ecker, C. M. Seifert, N. Schwarz, and J. Cook, "Misinformation and its correction: Continued influence and successful debiasing," *Psychological Science in the Public Interest*, vol. 13, no. 3, pp. 106–131, 2012, doi: 10.1177/1529100612451018.

- [14] Blackbird.AI, "Indian voters inundated with deepfakes during the largest democratic exercise in the world," *Blackbird.AI Blog*, 2024. [Online]. Available: <https://blackbird.ai/blog/india-election-deepfakes/>
- [15] The Tribune, "1 in 4 Indians came across political content that turned out to be deepfake: Report," *The Tribune*, Apr. 25, 2024. [Online]. Available: <https://www.tribuneindia.com/news/science-technology/1-in-4-indians-came-across-political-content-that-turned-out-to-be-deepfake-report-614539>
- [16] Asia Pacific Foundation of Canada, "AI amplifies political reach but magnifies disinformation in India elections," *APF Canada*, 2024. [Online]. Available: <https://www.asiapacific.ca/publication/indian-election-use-of-ai-political-campaigns-voter-engagement>