

AI-Generated Fake News Affects People's Trust in the Technology and the Medium

Siqi Tang^{1*}, Yongru Chen¹, Ruoxin Liu²

¹Hong Kong Baptist University, Hong Kong, China

²Shanghai Qibao Dwight High School, Shanghai, China

**Corresponding Author. Email: siqitangtang@gmail.com*

Abstract. Artificial Intelligence-generated content (AIGC) has profoundly influenced journalism by automating news production and reshaping audience–media interactions. However, public skepticism toward AIGC raises new challenges to media credibility and technological trust. While previous studies have explored AIGC in marketing and consumer contexts, limited attention has been given to its reception in news environments. This study addresses this research gap by examining public attitudes toward AI-generated news through a machine-assisted thematic analysis (MDCOR) of comments collected from two YouTube news segments. Four dominant themes emerged: widespread distrust of artificial intelligence and mainstream media, concerns over misuse and regulation, debates about the detectability of synthetic content, and criticism of media bias. Research findings indicate that AIGC exacerbates existing distrust rather than fostering new trust. This study contributes to understanding how generative journalism affects public trust and offers practical insights into transparency, provenance verification, and media literacy enhancement.

Keywords: Artificial Intelligence, consumer trust, news, communication

1. Introduction

News has long played a fundamental role in human society as both an instrument of communication and a public record of events. Historically, news served as a mechanism for communities to access reliable information about politics, economics, social affairs, and international developments. From early printed newspapers in the seventeenth century to the rise of radio and television in the twentieth century, news evolved into a trusted medium that not only informed but also shaped public opinion and collective memory. Over time, the daily consumption of news became deeply embedded in people's routines, reinforcing the perception of news organizations as authoritative sources of knowledge.

The emergence of digital media has transformed the traditional news landscape. Although professional journalism still operates within established institutions, but the continuing rise of the internet and social media platforms such as Facebook, Twitter (X), and YouTube has democratized the production and circulation of news. Since this trend makes information spreads more rapidly than ever before, the way people prefer to get news shifts from printed newspapers to digital screens now. Also, social media in particular has become a main channel for public to obtain information,

which allows users instantly access to news. However, it blur the boundaries between professional journalism and AI-generated content, and public sometimes is hard to distinguish them.

At the same time, with the changes in the media industry, artificial intelligence (AI) technology has developed rapidly into one of the most effective technologies of the twenty-first century. For example, no matter autonomous driving in Tesla vehicles or advanced facial recognition systems, AI has significantly permeated in people's everyday life. Additionally, among many applications of AI, one of the most controversial usages is the use of AI to generate synthetic text, images, and videos. Especially in the field of journalism, AI-generated news has become as a common phenomenon, like in algorithms created or news reports presented. For example, the Associated Press has used AI to automatically produce financial summaries since 2014, and China's Xinhua News Agency generated a virtual AI news anchor to report the news. On social media platforms such as YouTube, AI-generated news channels attract substantial viewers, and the videos they published gained thousands of views, likes, and comments. The proliferation of such content shows AI's development influence on production, distribution, and consumption ways of news.

Although not doubt that the using AI to produce news production can improve efficiency and reduce businesses' costs, it also raises public concerns about regarding authenticity and credibility of the information. The possibility of AI-generated fake news lead to the relationship between audiences and the media has become weaker, because the reliability of information became more uncertain than before. For many, the blending of "real" and "synthetic" reporting triggers skepticism toward both the technology and the media institutions that adopt it. Questions arise: How do people perceive AI-generated news? Does it diminish trust in AI technology as a whole, or does it specifically undermine confidence in journalism?

To address this research gap, this study employs qualitative content analysis to explore public attitudes toward AI-generated news and its impact on trust in technology and media. Specifically, we analyze the comment sections of two widely viewed AI-generated news videos on YouTube (<https://www.youtube.com/watch?v=SRA4brHXBBQ> and <https://youtu.be/7B-J89xKL7o?si=sTzGHCDdQqySYZkt>), using the MDCOR data analysis tool to identify patterns in public discourse. The research focuses on three key dimensions: attitudes toward AI-generated news content, perceived trust in AI technology, and perceived credibility of media that use AI. By comparing and synthesizing these perceptions, this study aims to uncover the underlying factors influencing trust and provide insights into how the rise of AI-generated fake news is reshaping the relationship between the public, technology, and the media.

2. Literature review

2.1. The evolution of media and the centrality of social media

Human communication has evolved through pivotal stages: from verbal exchanges and inscriptions to the revolutionary inventions of the printing press and electronic media. Each stage has progressively expanded the reach and reduced the constraints of information dissemination [1]. The advent of the internet and digital media marked a paradigm shift, enabling the creation, distribution, and modification of content in multiple formats with unprecedented convenience and interactivity [2].

For contemporary organizations, operational proficiency in social media is no longer optional but a critical skill for navigating the digital age.

2.2. The rise of artificial intelligence in content generation

The field of Artificial Intelligence (AI), formally established at the 1956 Dartmouth Conference, has undergone significant evolution from its early theoretical foundations [3, 4]. While early generative technologies faced commercialization challenges, breakthroughs like the Transformer model [5] and subsequent large language models (LLMs) such as OpenAI's ChatGPT have dramatically enhanced AIGC capabilities, enabling its large-scale application across various domains [6, 7].

This technological progress has led to the widespread adoption of AI across numerous industries, including government, education, and healthcare, where it automates processes and improves efficiency [8, 9]. The media sector, particularly journalism, has been a notable adopter. News organizations now leverage AI across three primary areas:

(1) Automated News Production: AI systems autonomously generate data-driven reports, as seen with The Washington Post's Heliograph and The Associated Press's financial news, streamlining workflows from data gathering to narrative generation.

(2) News Delivery: The development of synthetic news anchors that operate 24/7, transcend language barriers, and minimize human error.

(3) Verification of false information: Utilizing machine learning, natural language processing, and graph-based technologies to quickly analyze content helps in identifying fake news.

2.3. Challenges, trust gaps, and the shift from consumer to public trust

Despite its promise, the application of AI in journalism faces multifaceted challenges that contribute to a potential crisis of trust, distinguishing it from the commercial context of AIGC.

(1) Functional and Reliability Challenges: AI tools struggle to verify the authenticity and accuracy of real-time information, sometimes producing "artificial hallucinations" or fabricated content [10]. They lack the ability to deeply analyze complex linguistic structures like irony and political metaphors, nor excel at investigative reporting requiring nuanced judgment [11].

(2) Ethical and Professional Issues: Data privacy breaches, copyright infringements, and opaque information sources erode public trust. Although AI news anchors are efficient, they often lack the emotional depth and interactivity of human journalists, leaving audiences feeling disconnected from reality [11].

(3) Inherent Systemic Flaws: AI's filtering logic tends to be simplistic and rigid. When rules become overly inflexible, this can lead to biased or incomplete news coverage [12].

Although the existing research has effectively addressed the issue of building consumer trust in AIGC in the business marketing environment, this research focus is insufficient for the field of news reporting. The credibility of news media is foundational to democratic society, and public trust in this institution is built on different pillars—such as perceived objectivity, accountability, and public service—compared to trust in a brand. The emergence of AI-generated fake news directly threatens this specific form of trust.

Therefore, this study addresses a critical void in the literature by shifting the focus from consumer trust in marketing to public trust in technology and the media within the news ecosystem. It aims to investigate how AI-generated news, particularly instances of misinformation, impacts people's trust in the AI technology itself and the credibility of the media institutions that employ it, a question that remains underexplored compared to the well-studied dynamics of consumer trust.

3. Data collection

While integrating AI into news production has enhanced operational efficiency for many media organizations—such as the Associated Press (AP) adopting AI in financial reporting since 2014 and Xinhua News Agency deploying AI virtual anchors—this technology has also sparked public skepticism about its reliability and the credibility of implementing it. The proliferation of AI-generated fake news in daily life poses a significant emerging risk that could substantially undermine public trust. This study examines evolving trust dynamics by analyzing public discourse surrounding AI-generated news. To capture authentic and direct public reactions, we selected YouTube user comments as the primary data source.

The selection of videos follows three key criteria: relevance to AI-generated news, high viewership and comment engagement, as well as timeliness. Two YouTube videos meet these standards: one released by NBC on July 26, 2025, and another by MSNBC on July 22, 2025. Video 1 (MSNBC) features news anchor Rachel Maddow exposing the proliferation of AI-generated misinformation on social media. Through fictional scenarios—including rescue operations for Texas flood victims, childbirth scenes, and founding a news network—she vividly demonstrates the dangers of deep fake technology. This clip aims to emphasize the importance of verifying news sources and viewing online videos with a critical mindset in today's information overload era. Video 2 (NBC News) explores how hyper-realistic AI-generated news videos are proliferating online, making fact-checking increasingly difficult. Experts warn that rapid technological development could allow unverified misinformation to spread widely, sparking new public concerns about digital information trust.

These two videos present nearly identical AI-generated news content in the same format, both released on YouTube through professional news media accounts and presented by news anchors via broadcast. As of August 1, 2025, NBC News' video had garnered 775,617 views, while MSNBC's video received 368,286 views. A total of 3,174 comments were collected from NBC's video and 1,899 from MSNBC's video, including 2,418 parent comments from NBC and 1,523 from MSNBC. Both videos demonstrated real-time relevance with closely timed publication. To ensure clarity and focus in analysis, all subcomments and replies were excluded, retaining only parent comments for qualitative content analysis using MDCOR tools. This methodology enables systematic exploration of public attitudes toward AI-generated news, trust perceptions of artificial intelligence technology, and judgments regarding media credibility.

4. Study design

After completing data collection, this study adopted a machine-assisted thematic analysis approach based on the Machine-Driven Classification of Open-Ended Responses (MDCOR) framework. MDCOR is an unsupervised, machine learning-based textual classification tool designed to efficiently and rigorously process large volumes of qualitative responses while retaining participants' original expressions [13]. It reduces researchers' subjective bias inherent in manual coding, enhances reproducibility, and offers scalability suitable for classifying thousands to tens of thousands of social media comments.

In this study, MDCOR was employed to identify latent thematic patterns from a large YouTube comment corpus while preserving authentic participant feedback—, establishing a foundation for subsequent explanatory and theory-driven analyses. All collected parent comments were imported into MDCOR in CSV format with unique identifiers and designated columns for text content. The software then performed automated natural language processing and text cleaning, including word

form normalization and removal of sparse terms or abnormal feedback below minimum frequency thresholds (e.g., terms appearing less than 1% of the time). Based on initial term frequency reports, high-frequency but low-distinctiveness words (e.g., "AI", "news") were selectively removed manually to optimize and refine the classification model.

Following text preparation, Bayesian sampling parameters—including burn-in iterations, Gibbs sampling draws, and hyperparameters α and β —were set with recommended defaults to maintain methodological consistency. MDCOR subsequently executed metric assessments to determine the optimal number of latent codes via convergence criteria. The resulting metric plot guided the selection of the final code quantity, balancing statistical optimality and thematic interpretability.

Once the model was trained, MDCOR output fully classified datasets, including:

- A document-level assignment of each comment to its most probable thematic code;
- Relative text contribution and group fit indices indicating the strength and consistency of classification;
- Interactive visualizations (e.g., intertopic distance maps and word frequency distributions) to support qualitative sense-making;
- Optionally, group-based comparative analysis (e.g., by commenter attributes) via chi-squared tests and network correlation plots.

This machine-driven process allowed for efficient, replicable categorization of public reactions to AI-generated news, while the retention of original comment text ensured that nuanced, context-rich interpretations remained accessible for deeper qualitative inquiry.

5. Result

5.1. Data preprocessing and text mining outcomes

This study took parent comments from two YouTube videos (released by MSNBC on July 22, 2025, and NBC on July 26, 2025) as the research object. These two videos both focus on AI-generated news, with view counts of 368,286 and 775,617 respectively as of August 1, 2025. A total of 4,341 parent comments were collected initially, including 1,523 from MSNBC and 2,818 from NBC. Subcomments and replies were excluded in advance to ensure the independence and clarity of the analysis objects, which is consistent with MDCOR's requirement for the integrity of single text units in qualitative analysis.

5.2. Determination of optimal latent codes (metrics assessment)

The first step is using MDCOR to evaluate the optimal number of codes by calculating four classic metrics [14-17]. The output shows that two key metrics converged on 4 codes as the optimal solution. The deleting metric proposed by Cao et al. minimizes term overlap between codes, and successfully reached its minimum value when the number of codes was 4; the metric proposed by Deveaud et al. maximizes differences between codes, and make sure the result reached its maximum value when the codes was 4 [15, 16]. According to the criterion proposed by Nikita & Chaney that "convergence of at least two metrics is sufficient for decision-making", and considering the interpretability of thematic meanings (4 codes can clearly distinguish different dimensions of public attitudes without excessive fragmentation), the final number of latent codes was determined to be 4 [18].

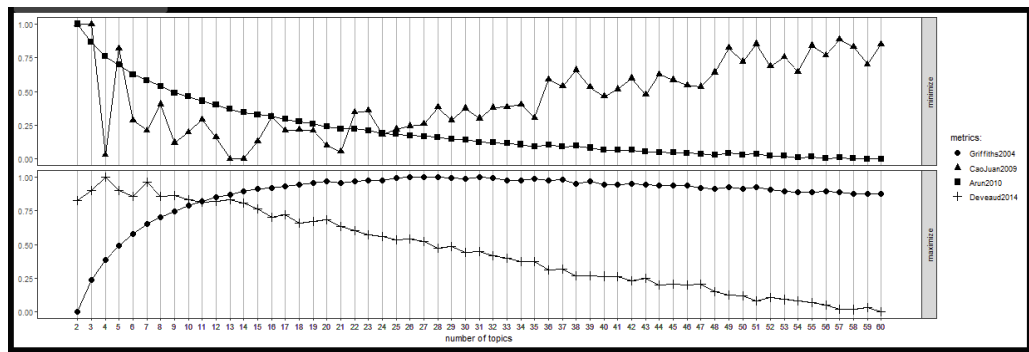


Figure 1. Metric convergence and code quantity selection

5.3. MDCOR classification results and thematic interpretation

Based on the above results, the contextualized meanings of most representative responses are shown below:

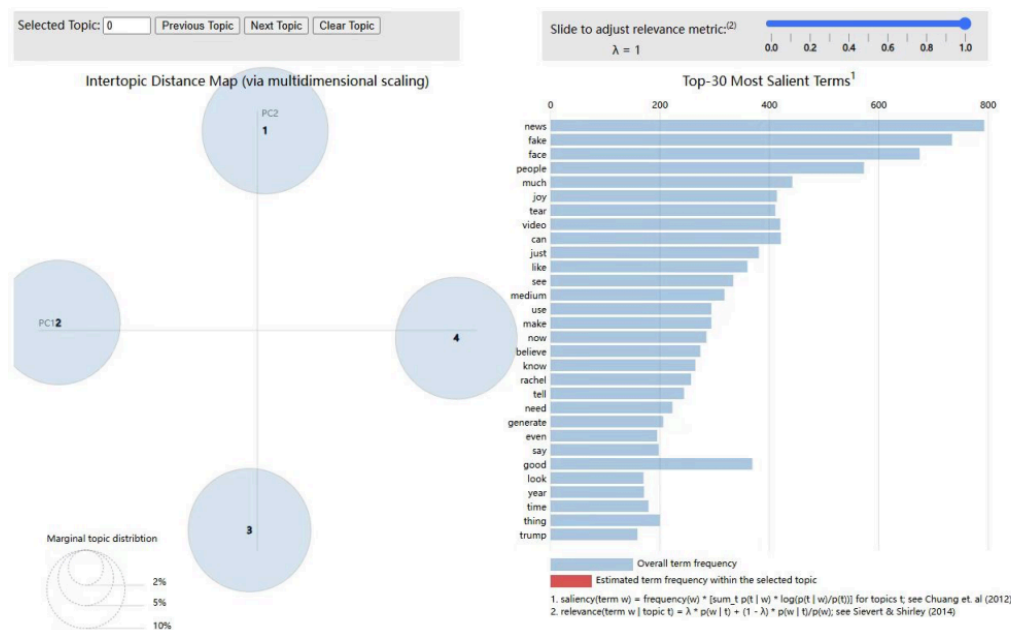


Figure 2. Interactive visualization of code-term relationships

5.3.1. Code 1 : deep distrust and criticism of media and AI

Category V1 comments focuses on topics related to the impact of technology on public thinking and media manipulation, which expressed strong distrust of the mainstream media, which is perpetuated by bias, manipulation and false reporting. The most representative comment is "It used to be you assumed people were telling the truth, until they were caught in a lie. Since trump, you assume everything is a lie until proven true. Now AI is spinning everything out of reality flushed face" It point out that AI-generated fake content is an extension of traditional media problems rather than a new threat. The second most representative comment is "Once again, the lack of critical thinking and digital literacy in our society is proving deeply harmful. Many Americans already struggle to discern reality, and without a major overhaul of our school curricula across all grade levels, this problem is only going to get worse." It shows the importance of the critical thinking skills and

specific technologies will not make much difference for people who are critical thinkers themselves. V1 comments shows that it's not the technology itself (such as artificial intelligence) that is really worrying, but the fact that some people are highly susceptible when they don't question the information they receive. Also, AI could be used by those in power to cover up the truth, falsify history, or evade responsibility, further weakening the public's judgment of information.

5.3.2. Code 2 : concerns about AI technology misuse and regulatory calls

Category V2 comments focuses on expressing deep concerns about the societal risks that AI-generated content may pose, particularly in legal, ethical, and public safety aspects. AI videos could be used to falsify evidence, frame innocents, manipulate public opinion, and even influence election results. Also, putting forward some suggestions to reduce the occurrence of risk problems.

The most representative comment is "AI videos could be used to falsify evidence, frame innocents, manipulate public opinion, and even influence election results." It call for international or national regulatory mechanisms, such as adding spectator marks, digital signatures to distinguish between real and AI-generated content.

The second most representative comment is "There was a very short period of time, from the popularization of smartphones until now, where common citizens could use videos as proof of things happening around them. With the advance of AI, that's ended. Videos can't have any credibility anymore. We'll need once more rely on big agencies and tv channels to watch true things. Which is a step back. And worse than that, the TV channels also get used to simply reproduce video that comes to them, so they'll need to readapt to a world where a video now has just a few percentage of chance of being real. There's also an even worse stuff, at some point criminals will be able to use videos and photos in social media to create perfect impersonations of common people. That opens a huge sphere of concerns. It's easy to say that a famous person didn't do some stuff, that their video is fake. But what about common people? Fake videos can cause huge damage to innocent people. It's simply not safe keep our videos and photos on the internet anymore. Never was, but with AI is worse." It points out the threats and dangers brought to people by the rapid development of AI, and the further development of videos and pictures should be suspended until people adapt to AI technology.

V2 comments reflect the lack of regulation in the development of artificial intelligence technology and the public's dissatisfaction with the adaptation process.

5.3.3. Code 3 : AI discriminability and universality

Category V3 comments focus on the discussion and reminder of people's AI recognition capabilities, and the generalization of AI fraud. The comments show that some people believe that AI counterfeiting technology is not perfect, showing their neglect of AI. What really triggers anxiety is the "democratization" of disinformation - anyone can fake it, not just the traditional big media or those in power. The most representative comment is "Here's how you can tell: Real people variate their tone, cadence, and facial expressions far more, and in a far less structured way. Lot of AI content has less than an octave of variation in tone, and sounds almost monotone. This is why you need to force people to not read stuff like a robot. The more and more repetitive and monotonous people's speaking gets the harder it is to tell. Way too many people on the internet talk with an annoyingly repetitive cadence and tonal pattern like they're an actual robot." It shows that AI-generated content looks more and more realistic, but it still lacks personality, voice, and expression match. Young people and experienced viewers can mostly tell the difference. The second most representative comment is "So, people who lie the old fashioned way, like the NY Times... Got it."

It proposes that in the past, only big media could make up credible lies, but now technology allows ordinary people to create content that looks real, which is called "democratization". V3 comments did not focus on the impact of AI fraud, but rather whether people can identify it as AI and discuss its ubiquity in today's life

5.3.4. Code 4: irony and double-standard accusations against the mainstream media

Category V4 comments lies a question of the credibility of traditional media and the society situation of the U.S, and the focus of the comments is on the critique of traditional media that spreading false information through omission and misdirection. The most representative comment is "oh and you legacy news is so REAL with your Life size Cardboard cutout of your ancors and your FAKE NEWS is REAL? NOPE you big News Networks been running FAKE News for over 50 years in your Big Studios and your Fake Camera Shots and your Edits and now with Fake Stories and Fake News for the Top News Networks in U.S.A. is is Very Sad you Calling other Channels FAKE LOL LOOKS WHO TALKING." It strongly accused mainstream media outlets of long-term biased reporting and propaganda, and arguing that they are not qualified to criticize AI-generated "fake news." This statement directly shows distrust of the authenticity of traditional media reports, believing that the media has serious biases and injustices in the process of information dissemination. It reflects public dissatisfaction and concerns about media behavior, emphasizing the need for the public to remain vigilant and think critically when receiving media information. However, the second most representative comment is "That's why the United States is rapidly going down the drain." This comment shows that over half of American youths do the stupid thing, which is relying on social media for everything, even news. Therefore, news outlets have to adapt and start getting in on the act of drawing in the views by entertaining. Also, since news is hardly ever entertaining, many online organizations resort to misinform and dramatize news. All of these actions aim to attract youth, so they can at least compete with social media in entertaining. However, the false actions also cause a wide spread of fake news. This theme focus on criticism to audiences who use social media to get news, which reflects how today's social media effect on the traditional industry, and shows some people's concern about this serious phenomenon and their understanding of how traditional media responds. The V4 commentary shows a different perspective on the spread of fake news by traditional media.

5.4. Association between codes and video sources

To explore whether there were differences in public attitudes (reflected by code classification) corresponding to the two video sources (MSNBC and NBC), the study used MDCOR's optional group comparison function to conduct a chi-square test of independence.

5.5. Chi-square test results

The red connecting line between the two themes signifies the maximum correlation. A hypothesis test will be displayed below. The color in the chi-square distribution chart 16 denotes the strength of the association: red stands for negative correlation, which implies that the software's hypothesis is higher than the actual observation; blue represents positive correlation, indicating that the software's hypothesis is less accurate than the actual observation. Dotted lines and rectangles have no statistical correlation implications. Overall, MDCOR adopts machine-driven approaches to identify and analyze the public's responses to the event, aligning with the study's objectives. It notably cuts down

the time costs related to manual classification and greatly supports thematic analysis. Additionally, this offers precise validation for the experimenter's thematic classification.

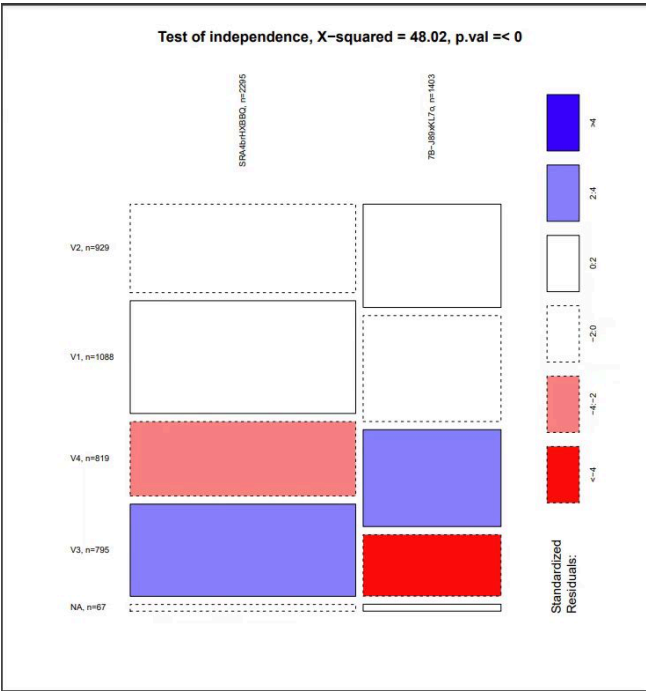


Figure 3. Chi-Square test result

6. Conclusion

From the introduction of the Turing Test in 1950 to the emergence of ChatGPT in 2022, over the span of 72 years, AI technology has evolved from concept to consumer application. While the advancement of artificial intelligence presents opportunities for journalism, it also poses challenges. How to appropriately view and utilize AI has become a critical issue that future journalists must contemplate and resolve in their relationship with artificial intelligence. This study emphasizes the impact of AI technology usage in the news industry on audience trust. This study explored public attitudes toward AI-generated news and its implications for trust in both technology and media, using a machine-assisted thematic analysis (MDCOR) of 4,341 parent-level comments collected from two high-viewership YouTube segments (MSNBC, July 22, 2025; NBC, July 26, 2025). The analysis identified four dominant themes: (1) pervasive distrust and sharp criticism directed at both AI systems and mainstream media; (2) heightened concern about malicious misuse of AI and calls for regulatory safeguards; (3) debate over the current discriminatory of AI-generated content and the uneven ability of audiences to detect fakes; and (4) accusations of media double standards, with legacy outlets themselves held responsible for undermining credibility. These themes indicate that AI-generated news exacerbates people's concerns about media bias and the lack of trust, while also bringing about new risks in terms of technology and ethics.

The findings have clear practical significance. To prevent abuse and rebuild trust, it is necessary to have a technical source, a more powerful platform, a legal framework, and maintain transparency within the news editorial department when using generating tools. Equally important is the continuous investment in public media literacy and critical thinking education, so that the audience can better evaluate digital content.

This study also has limitations. It only focused on two YouTube videos, only included parent comments, and used a single time window - this limited the ability to draw broad generalizations. Future research should increase the sample size across platforms and cultures. In conclusion, in the face of the new opportunities brought by artificial intelligence to the journalism industry, journalists must actively adapt to its technological logic. By mastering the methods and tools of artificial intelligence, they can better participate in and serve news production. Critical thinking and profound insights represent humanity's core competitive edge against AI. Simultaneously, we must confront the information crises AI brings, such as information overload and misinformation. The emergence of information crises is not merely a challenge for journalism but a societal issue affecting governments, the public, and all stakeholders. Robust information protection laws, a well-informed and discerning public, and comprehensive media oversight are essential pathways to safeguarding information value in this era of information crises.

Acknowledgement

Siqi Tang, Yongru Chen, and Ruoxin Liu contributed equally to this work and should be considered co-first authors.

References

- [1] Ukpe, A. J., & Bassey, E. Y. E. (2023). *The History of the Print Media. Akwa-Cross People of Nigeria: History, Heritage, and Culture*, 287.
- [2] Storage, T. E. (2013). *Technology brief. IEA-ETSAP and IRENA© Technology Brief E17-January*
- [3] Crevier, Daniel, *AI: The Tumultuous Search for Artificial Intelligence*, New York, NY: BasicBooks, 1993, ISBN 0-465-02997-3
- [4] Russell, Stuart J.; Norvig, Peter, *Artificial Intelligence: A Modern Approach 2nd*, Upper Saddle River, New Jersey: Prentice Hall, 2003, ISBN 0-13-790395-2.
- [5] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30.
- [6] Ray, P. P. (2023). ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. *Internet of Things and Cyber-Physical Systems*, 3, 121-154.
- [7] Wang, Y., Pan, Y., Yan, M., Su, Z., & Luan, T. H. (2023). A survey on ChatGPT: AI-generated contents, challenges, and solutions. *IEEE Open Journal of the Computer Society*, 4, 280-302.
- [8] Fui-Hoon Nah, F., Zheng, R., Cai, J., Siau, K., & Chen, L. (2023). Generative AI and ChatGPT: Applications, challenges, and AI-human collaboration. *Journal of information technology case and application research*, 25(3), 277-304.
- [9] Lao, Y., & You, Y. (2024). Unraveling generative AI in BBC News: application, impact, literacy and governance. *Transforming Government: People, Process and Policy*.
- [10] Gutiérrez-Caneda, B., Vázquez-Herrero, J., & López-García, X. (2023). AI application in journalism: ChatGPT and the uses and risks of an emergent technology. *Profesional de la Información*, 32(5).
- [11] Cardaş-Răduța, D. L. (2024). The effectiveness and limitations of artificial intelligence in journalism. *Saeculum*, 57(1), 111-119.
- [12] Eckroth, J., Dong, L., Smith, R. G., & Buchanan, B. G. (2012). NewsFinder: Automating an AI news service. *AI Magazine*, 33(2), 43-43.
- [13] Canché, M. S. G. (2023). Machine driven classification of open-ended responses (MDCOR): An analytic framework and no-code, free software application to classify longitudinal and cross-sectional text responses in survey and social media research. *Expert systems with applications*, 215, 119265.
- [14] Arun, R., Suresh, V., Veni Madhavan, C. E., & Narasimha Murthy, M. N. (2010). On Finding the Natural Number of Topics with Latent Dirichlet Allocation: Some Observations. In M. J. Zaki, J. X. Yu, B. Ravindran, & V. Pudi (Eds.), *Advances in Knowledge Discovery and Data Mining* (pp. 391–402). Berlin Heidelberg: Springer.
- [15] Cao, J., Xia, T., Li, J., Zhang, Y., & Tang, S. (2009). A density-based method for adaptive lda model selection. *Neurocomputing*, 72(7–9), 1775–1781. <https://doi.org/10.1016/j.neucom.2008.06.011>

- [16] Deveaud, R., SanJuan, E., & Belloto, P. (2014). Accurate and effective latent concept modeling for ad hoc information retrieval. *Document numérique*, 17(1), 61–84. <https://doi.org/10.3166/dn.17.1.61-84>
- [17] Griffiths, T. L., & Steyvers, M. (2004). Finding scientific topics. *Proceedings of the National academy of Sciences*, 101(suppl 1), 5228–5235. <https://www.pnas.org/content/101/suppl1/5228>.
- [18] Nikita, M., & Chaney, N. (2020). ldatuning: Tuning of the Latent Dirichlet Allocation Models Parameters. R Project. <https://CRAN.R-project.org/package=ldatuning>.