



AI-Generated False Media and Trust in the Media

Siena Khym

1. Glen Rock High School

Abstract: Media deceit spreads through social media, mainstream news, and published articles. Artificial intelligence and misinformation can now convincingly mimic reality through technology. This blurs the line between fact and fiction, threatening to divide society. Misinformation influences voting behavior on education, healthcare, and the economy. The study examines the impact of realistic false media on media trust when individuals must distinguish between AI-generated and human-authored texts. I believe AI-generated writing and identification will harm journalism. Thirty-two participants, aged 13 to 80, were surveyed in four sections after demographic inquiries. Initially, participants identified AI-generated and human-authored materials. In the second session, participants distinguished between human-authored texts and AI-generated texts. They were asked if human and AI articles represented the same event. A follow-up questionnaire examined the erosion of faith in the media. It's hypothesized that the participants will be more skeptical of the media after exposure to AI-generated text. Media skepticism will likely increase with age, particularly among individuals who have difficulty recognizing AI-generated content. Deceptive media erodes faith in journalism by impairing individuals' ability to discern fact from fiction. As media credibility erodes, individuals are increasingly forced to make high-stakes judgments based on unreliable information, accelerating social polarization.

Keywords: Artificial Intelligence, Misinformation, Media Trust, AI-Generated Content, Fake News, Media Literacy, Source Credibility

INTRODUCTION

People have been exposed to misinformation through the media since the 1800s, but today, a more pressing concern is emerging: artificial intelligence. It was even reported in 1835, when articles from the New York Sun claimed that life on the Moon had been discovered. Later, after they gained more subscribers to the newspaper, they admitted that the story was fake. This example reveals that misinformation can serve economic interests when spread in the news, as the New York Sun created this news to boost sales and make a profit. Then, in 1938, Orson Welles' War of the Worlds radio broadcast led many people to believe an alien invasion was underway, which caused widespread fear among Americans. Another instance in 2003, during the lead-up to the Iraq war, articles were spread globally in newspapers about Saddam Hussein's weapons of mass destruction. These weapons, however, did not exist. Both of these examples created immense fear in readers. The newspaper put out fake news that made many people very panicked that something horrible was going to happen, like an alien invasion or an attack from deadly weapons. The readers had a lot of trust in the newspaper, and what they saw was what they believed to be true. Now, after discovering that the newspapers' reports were untrue, it became less likely for people to trust them. This could lead to a serious global event that puts the reader at risk; they won't believe the newspapers, since they haven't been truthful in the past. The spread of false information isn't a new phenomenon in the media.

Although fake news has been a part of the media for centuries, what is new is the advent of Artificial Intelligence (AI). AI enables fake news to be more realistic, which raises a more serious issue. It is now harder to distinguish between what is real and what is fake due to the rise of AI. With AI, videos, images, audio, and text can be produced that look real and are easy to generate. According to the New York Times, there was an instance in 2025 where President Donald Trump posted on social media a video that was generated by AI of former President Barack Obama getting arrested in the Oval Office by the FBI. Then a clip of Obama in a jail cell is shown. This is deeply concerning since the video could have led people to believe that the former president got arrested, which could have serious repercussions for the nation.

The purpose of this study is to understand how well people in today's society can determine whether a text is generated by AI or if it was written by a human. The study investigates people's ability to determine whether a text is AI-generated, to distinguish AI-generated text from human-written text when presented together, and to assess whether information from an AI text or a human text is more believable to humans. With false media being so prevalent, whether it be on news outlets, radio shows, podcasts, and social media, more and more people are being exposed to false information. While other studies focused on misinformation, this study focuses on trust in the media. The study aims to see if people are losing their trust in the media due to being introduced to realistic false press. AI is a tool that is relatively new, and the impacts of AI on people continue to be studied.

Social media is also a source of information for many people. On these platforms, AI and fake news can be rapidly, cheaply, and easily spread and reach a wider audience than ever before. A study from the Pew Research Center aimed to understand who gets their news from social media. They found that around one in five Americans get their news from influencers on social media regularly, and this number is higher for adults under the age of 30, with 37% of adults. It was also observed that more news influencers were explicitly identified as Republicans (27% of news influencers) than as Democrats (21%). Meaning that the news that people get from social media influencers might be biased towards one political side. Additionally, 63% of news influencers are men, which can also provide bias from a perspective of only men. What is most concerning is that the majority of the news influencers don't have any background with a news organization. 77% of these influencers don't have experience in journalism. This is an epistemic crisis, where these people just don't have enough data and information to be able to have the knowledge to spread information to others. This can be dangerous because now the news is transitioning from being shared by trained professionals to uncredentialed individuals, and this could lead to inaccurate information being spread and believed as true. With the rise of people getting their information from social media, the extent to which fake news can be exposed to people is greater than ever before, due to the unqualified news influencers spreading the news and the ability to spread information to a vast number of people in a short amount of time for cheap.

There have been studies on AI-generated fake news, and this research has provided people with important information on its effects. GPT-2 was used to generate AI articles, and researchers collected both real and fake human-written articles for a study. All the AI and real were about the COVID-19 virus. They then input the articles into detectors that were designed to distinguish between fake and real articles. It was found that the detectors performed better at detecting what was fake in human-written articles than in AI-generated

articles (Diana Trandabăța et al., 2023). Some of the articles made by AI were classified as real, which shows that AI-generated false news can be as convincing as false news created by humans.

When looking at Facebook posts about fake news and the comments on those posts, the researchers observed that the main features of the news, the content creator, and comments explained why people engage with the fake news. Pictures or videos sparked more of a reaction from viewers when the news triggers sadness, anger, or fear. Additionally, headlines that were created for clickbait with multimedia got more engagement, and articles that didn't match the subject of the comments or other news got less engagement (Neha Chaudhuri et. al. 2025). This shows that fake news can create a reaction from viewers and reach massive engagement through the use of social media.

In addition to observing how false news spreads on social media, other studies have examined why people share fake news on platforms, specifically Facebook users, about the COVID-19 virus. What was observed was that people who had a higher education reported a higher literary skill, and individuals with high literary skills shared fake news less. Additionally, people who were more trusting shared fake news more. If the motivation for someone sharing information was for extrinsic reasons (political beliefs, religion, image, etc), the more fake news they shared than people whose motivation for sharing was intrinsic (personal reasons like fun, socializing, or altruism) (Christaine Melchoir et. al. 2025). This shows that false media is shared by social media for a variety of reasons, and trust can have an impact on the spread of false media as well.

Additionally, a study examined whether people believed misinformation more if they had close ties or weak ties with the person sharing it. Participants were shown four Facebook-style headlines. Participants were shown Facebook-style headlines, and all were shown all four combinations of the conditions of tie strength and truthfulness. From that, it was found that strong ties are more trustworthy than weak ties in terms of ability, integrity, and benevolence. However, false stories from weak ties were more believed than false stories from strong ties (Babajide Ostayi et. al. 2024). This shows that the source of the information does matter, and it impacts whether we believe the information or not. It also suggests that people may be more likely to believe a source that they don't know that well, like a person on the news or an influencer on social media, over someone close to them that they know personally.

In addition to studies about how the strength of the relationship of the source impacts believability, the impact of the source's political opinions has also been studied. Both conservative and liberal participants looked at Twitter headlines from both liberals and conservatives, and it was found that people through news sources that didn't share their political beliefs were more biased, with liberals rating conservative sources as more extreme and conservatives rating liberal sources as more extreme. Moreover, people trusted and judged news from sources that shared their political views more accurately. Liberals judged liberal misinformation as more accurate than conservative, and conservatives judged conservative misinformation as more accurate than liberal. It was also found that sources from liberal to conservative perspectives were less credible, and sources from conservative to liberal perspectives were less credible. Along those lines, liberals trust liberal sources more, and conservatives trust conservative sources more (Cecilie Steenbuch Traberg et. al. 2022). From this information, it can be concluded that your political view and the source's

political view can impact whether you believe the information or not. If a source aligns with your political view, you are more likely to believe in the misinformation than if the source doesn't align. This can have major consequences because if you believe misinformation about politics just because they share your political opinion, it could impact your vote in an election and your other political views as well.

Other studies have focused on misinformation and politics was a study conducted as well, which aimed to understand how personality traits impact whether people believe political misinformation as true and whether they would share the information. Participants were shown a series of false news headlines from social media right before the 2020 election. Like the other study done by Cecilie Steenbuch Traberg et. al, it was found that Republicans were more likely to believe conservative misinformation and Democrats were more likely to believe liberal misinformation. They also found that older people were less likely to believe and share political misinformation. In regard to how personality plays a role in believing the misinformation, it was found that more agreeable people were better at spotting false conservative misinformation as false and people who are more extraverted as more likely to believe pro-conservative misinformation. Additionally, the people who were less likely to share misinformation were people with high agreeableness or conscientiousness. Another finding of this study was that people with higher cognitive ability were better at identifying misinformation and less likely to share it. It was also observed that personality traits had more impact on people with low cognitive ability, while for people with high cognitive ability, personality had little effect on sharing or spotting misinformation (Saifuddin Ahmed et. al. 2022). This study conveys that personality and cognitive ability do affect whether we believe misinformation or not, and whether we share the misinformation. Understanding our own personality traits can help us.

Since the AI faces were becoming more realistic and getting increasingly harder to tell if a face is real or fake, there was a study aimed to determine the factors that influence whether people think a face is real or fake. After showing a series of faces, it was observed that 44% of the faces were judged as fake and 56% as real, even though the faces were all real. When examining the determinants of reality judgments, more attractive males and females were judged as more real. Additionally, higher trustworthiness in female faces increased the likelihood of judging a face as real. Higher familiarity with male faces increased confidence in real judgments and decreased confidence in fake judgments (Dominique Makowski et. al. 2025). This study shows that gender and facial features can affect whether a person believes something is real. It can also be concluded that if someone believes a stimulus is potentially fake, it might lead them to mistakenly believe something real is fake. So, with the increasing amount of misinformation and fake news in the media and people having knowledge of this, it could lead to people starting to not believe in true information they find in the media, and instead think that real information is fake.

Many previous studies on fake news and AI have focused on the overall impact of misinformation and the effects of AI on viewers. In addition, these past studies have primarily focused on fake news about the COVID-19 virus due to their vast amount of misinformation surrounding the disease in the news and on social media. These studies also focused heavily on social media's impact on the spread of fake news. While these studies do provide us with more understanding about the impact of fake news and AI, they don't focus on how it affects people's trust. This study also aims to explore fake news beyond COVID, since the impact of false media during that period has already been studied. This study aims

to find out if people believe that AI-generated texts are written by humans and if people trust the information from AI-generated texts.

There is epistemic uncertainty about AI and AI-generated false media, since it has only been around for a couple of years. AI is constantly evolving, and the issue of AI becoming realistic enough to be hard to decipher if it is real or not has only started to become a problem within the past year. So there is some uncertainty about what we know about AI and its impacts due to there not being much data on the subject yet. There is a gap in understanding AI in general, especially with the impact of AI on trust. Additionally, AI-generated content is a prime example of a collapse of the traditional truth-validation mechanism. Because AI is very new and not yet well studied, it is difficult to understand its nature. It is almost impossible to predict what AI is going to produce and, therefore, extremely difficult to know if something is true or false. Nothing actually is fully true or fully false due to the lack of research data on AI. AI is a very new topic, and due to that, there is a lack of knowledge and understanding of AI.

The spread of false news and people's growing distrust of the media can have serious consequences. Today, false media is becoming very realistic through AI, which can lead to more serious outcomes than ever before. The loss of trust can impact public health, democracy, and society. People start to believe in untrue information, which can lead them to lose confidence in other organizations. People might not trust the healthcare system, leading them to forgo medical care they need because of misinformation they see in the news. Much of the misinformation around healthcare and other important topics are spread through politicians and other influential people, which could have major effects on democracy. When people start to believe that fake events happened or any false information about the government or spread by the government, it affects their decision of who they vote for in an election. This could lead to a government official being elected who isn't fit for the role, with major effects on democracy. Due to the consequences for democracy, society as a whole can be affected as well. False information can cause society to become divided and further polarize an already existing divide. Realistic false media is a pressing issue with significant impacts in today's world that need to be addressed. AI-generated false media can be created much more cheaply and easily today, and when people are exposed to the false media, there can be major impacts on a person's trust. This study seeks to enhance understanding of the broader implications of misinformation, especially in the news. This pertains to factors such as democratic resilience and public well-being.

When people see realistic fake media, it blurs the line between reality and fake. People start to be unable to tell the difference between the two. When they see a text that is fake and generated by AI, it doesn't just affect their trust in that one text; it might affect their trust in other texts as well. In the future, when they read texts, they might be unsure whether that text is fake or real. This can cause distrust in the media, as people struggle to fully believe what they see in the news.

LITERATURE REVIEW

A study by Manuel Goyanes et al. examined how AI literacy and attitudes towards AI affect people's trust in AI-generated news over time. The study included two survey waves: one in November 2024 and the other in May 2025. This resulted in 1,815 participants in the United States, 1,811 in Spain, and 1,802 in Chile. AI literacy was measured by an individual's basic

comprehension, recognition skills, functional application, and perceived performance of AI. Artificial intelligence attitudes were measured as individuals' general attitudes toward AI, including emotional responses, trust, integrity, and intentions. Their trust in AI news was measured by how individuals perceived AI-generated news as trustworthy and accurate. The study didn't conclude a positive relationship between literacy and trust. However, the study found a positive relationship between attitudes and trust. These results were consistent across the three countries. This shows that if someone views AI more positively, they are more likely to trust AI-generated news. If one has a more positive attitude towards AI, they may already be more trusting of the accuracy of the information in AI-generated texts, making them more likely to trust the AI-generated news. Additionally, those who have a poorer attitude towards AI may ultimately regard the information from AI-generated news as false. People's preexisting biases often affect their decisions, opinions, and trust.

Many consumer products are advertised in the media, which requires potential consumers to decide whether they trust the advertisement to be factual and an accurate representation of the product. AI technologies can produce realistic advertisements in media that may affect people's trust. A study by Rajesh Kumar Srivastava et al. examined in greater depth how AI-generated content affects trust by investigating the influence of AI-enabled technologies on consumer responses, such as trust. To gather information on this topic, researchers administered an online survey to 18-to 50-year-olds familiar with AI technologies. There were a total of 130 responses. The survey contained four sections focused on AI's influence on consumer decision-making, consumer trust and perceptions, and the ethical implications of using AI in social media marketing. It was found that the use of AI to personalize marketing enhanced consumer trust through relevance, efficiency, and accuracy. Ads that are more personalized are more effective at convincing consumers to buy by making them feel that this product would be a perfect fit for their lives. The advertisement has to gain the trust of the people that the product is good enough to spend their hard-earned money on, and they do this by making the advertisement more accurate to what the people want in a product, more relevant to the culture of today's world, and more efficient at getting the information across. However, artificial intelligence doesn't always increase trust. Viewers only trust it when they perceive its responsible use. It was found that transparency and ethical governance are essential to earning consumer trust. So, to gain the public's trust, companies should ensure they use AI to make their marketing the best representation of the actual product sold. The more accurate the marketing is to the product, the more it will create a reputation that the company is reliable. If a company sells what it says, it will be trusted.

When using social media, people often seek or encounter medical advice on those platforms. While some of this content can be harmless, much of the medical advice, especially that created by AI, is inaccurate and spreads to viewers false information that can be detrimental to their health. A study by Shruthi Shekar et al. examined how AI-generated medical responses are perceived and evaluated by non-experts, who are the majority of people on social media seeking medical advice. To do this, researchers have a total of 300 participants provide evaluations for medical responses, either written by a medical doctor or by generative AI. Physicians labeled the AI-generated content as either high-accuracy or low-accuracy. The results of this research show that participants could not effectively distinguish between AI-generated responses and those of doctors. The participants rated high-accuracy AI-generated responses as significantly more trustworthy

than the doctor's responses, and rated low-accuracy AI-generated responses as similar to the doctor's responses. The average person looking on their social media for answers about some medical issue they are having or advice on how to live a healthy life would more likely trust responses that are generated from a technology program that is unqualified and error-prone, rather than from a person who went to school, got a degree, and gained experience with real patients. Doctors are trained to provide patients with accurate medical advice that will address whatever medical issue they are experiencing, whereas AI relies solely on statistics from humans. AI doesn't have any personal experience with patients that doctors have, which is critical for doctors to provide more accurate medical advice. Since people cannot distinguish between the doctor's responses and AI-generated ones, this may lead participants to trust AI-generated medical advice, which can be inaccurate. In the study, participants found low-accuracy AI-generated responses to be trustworthy. This shows that trust can form even if accuracy is weak. The authorship of the article may impact their trust, not just the content itself. They also indicated a high likelihood of following the inaccurate medical advice. The false alarm of seeking unnecessary medical attention isn't very detrimental to the viewer; however, what can be detrimental is trusting medical advice that prompts participants to take actions harmful to their bodies and not approved by doctors. Many health content online promotes diets that may result in weight loss, but cause you to lack the necessary nutrients our body needs, or to take certain supplements or eat certain foods that have no proven health benefits, or the opposite, or scare people into not eating certain foods that have no proven health drawbacks. In the study, the participants thought the AI-generated responses were more thorough and accurate. Even though they may sound accurate, the non-professionals, that may not always be the case. Believing false information spread by AI can have major implications for our health, especially when it is taken more seriously than advice from doctors.

HYPOTHESES

I believe that after being exposed to AI-generated texts and having to discern between human-written and AI-generated texts, people will lose faith in the credibility of news organizations. I believe it will be difficult for participants to differentiate between the two types of texts and to realize how similar AI-generated texts are to human-written texts, which may lead them to think that AI could generate the texts they see in the media. This may lead them not to trust the media, given their realization of how realistic AI-generated texts are. The study aims to determine whether exposure to AI-generated texts and the need to distinguish between human-written and AI-generated texts will lead to a loss of trust in the media.

I additionally hypothesize that participants over the age of 30 would be more likely to lose their trust in the media. People over the age of 30 aren't typically on social media as often as those under 30. The younger generation is more exposed to AI-generated content than the older generation, as social media is where most of it is found. This gives people under the age of 30 more experience in determining if a text is AI-generated or human-written. Moreover, people over the age of 30 have been participating in the media for years. They aren't used to having to discern between AI-generated texts and human-written texts, and remember a time when this wasn't an issue. Younger people have less experience with the media, and, especially for teenagers, this has always been a problem. They are more

used to having this problem when engaging in the media due to not having as much experience where this wasn't an issue. For the same reasons, I predict that the participants who engage in the media less frequently will be more likely to lose their trust in the media. This is because they have less experience deciphering between AI-generated texts and human-written texts, and aren't exposed to the media as much compared to those who frequently engage with the media. Since they don't have the experience in distinguishing AI-generated texts from human-written texts, they don't have a lot of prior experience to base their answers on. This may make the activities in the study harder, which may lower their confidence in being able to detect AI-generated texts. Having lower confidence in their ability may lead to lower trust in news outlets, because if they can't trust their own abilities to distinguish between AI-generated texts when engaging the media, they wouldn't be able to trust that what they are reading is certainly written by a human.

I also predict that those with poorer perceived ability to identify AI-generated texts would be more likely to lose trust. Participants aren't informed of their performance. This created an uncertainty that may drive lower trust because they don't know if they were proficient at detecting AI-generated texts or not. Those with lower confidence in their abilities may perceive that they performed poorly in detecting which text is AI and which is human. This may lead them to believe they aren't proficient at this skill to a greater extent than before the study. This may lead them to distrust their own judgment about a text's authorship when using the media. They may have lower trust in the media due to not trusting their own abilities to determine what is real versus what is fake. Participants who are more confident in their abilities are less likely to doubt their decision on whether a text is AI. This may result in them not losing trust as much due to their confidence in being able to determine the authorship of texts.

I also predict that participants who perform more proficiently at correctly identifying AI-generated texts will show a small decline in trust in the media compared to those who struggle to distinguish. Those who are more proficient will more likely trust the media because they have a better ability to know what is real versus what is fake. They won't have to worry that they will be manipulated into believing something is fake since they can determine whether AI is generated from human-written text. Those who are less proficient don't have these abilities. They are more likely to fall victim to believing in misinformation, which may lead them to lose trust in the media.

METHODS

There were a total of 32 participants in the study. The study population ranged from 13 to 80 years old. The majority of participants were teenagers, with 20 aged 13 to 19. The sample was mostly female, with 26 female participants. The rest were identified as male. The participants were administered an online survey to complete on their own. They were asked to take the survey on a laptop, with the form as the only open tab, in full screen, at medium brightness. They were also asked to place the computer screen parallel to their face on a table in front of them, about 2 feet away. Controlling the screen setup was important to ensure their ability to read the paragraphs clearly didn't affect their response. Additionally, they were asked to keep the room well-lit, quiet, and free of distractions or tasks while completing the survey. They were asked to complete the survey during the daytime, when they feel energized. Additionally, all texts involved in the study were either written by Siena

Khym or generated by ChatGPT using human-written texts. The same person created and adapted all texts to maintain consistency. This is because all human-written texts share the same writing style and skill, and all AI-generated texts were prompted to write in a style similar to the human-written text on a different topic.

The first section of the survey asked demographic data, including their age and gender. They were also asked in this section to rank their frequency of media use, their familiarity with AI, and their confidence in detecting AI-generated texts. A ranking of 1 indicated least frequent use, least familiar, and least confidence, and a ranking of 5 indicated most frequent use, most familiar, and most confidence. The next section of the study would present individual texts and ask participants to decide whether each was generated by AI or written by a human. Each text was 6-8 sentences long. There were 10 articles in total. 5 articles were AI-generated, and 5 articles were written by humans. After they read the text, they would indicate whether they thought it was AI-generated or human-written. Then, in 1-3 sentences had to provide their reasoning for their decisions and rank their confidence in the accuracy of their response on a scale of 1 to 5, 1 being the least confident and 5 being the most confident. The following section required participants to discern which text was generated by AI and which was written by a human when two texts were presented together. There were a total of 6 articles, 3 generated by AI and 3 written by a human. An AI-generated text was paired with a human-written text about the same topic. They were labeled Text 1 or Text 2. The participants then had to determine which of the texts was written by a human and which one was generated by AI, Text 1 or Text 2. They then had to provide their reasoning for choosing one text as AI-generated and the other as human-written. They had to rate their confidence in their answer on a scale of 1 to 5, with 1 being the least confident and 5 the most. In the next section, participants had to determine whether an event mentioned in a text actually occurred, using AI-generated and human-written texts. There were a total of four articles. One article was about a real event and written by a human, one article was about a real event and generated by AI, one article was about a fake event and written by a human, and one article was about a fake event and generated by AI. After they read each text, they had to identify if the event in the text actually occurred; they had to indicate their reasoning why they decided if the event happened or didn't happen, and they had to rank their confidence in the accuracy of their response on a scale of 1 to 5, 1 being the least confident and 5 being the most confident. The final section included open-ended questions to measure participants' lack of trust and other reflection questions. In 1-3 sentences, participants had to answer three questions: what did you use to decide whether it was AI or human, which passage was the hardest to determine if it was real or fake, and did this activity make you not trust the media in the future. Their response to the last question is how their loss of trust upon completion of the study was measured. However, this measure is limited since one open-ended question may not be adequate to capture people's trust in the media. After the participants completed the study, to measure longer-term changes in trust, about two weeks later, they were sent a secondary form that further observed how their trust in the media had changed since completing the survey.

RESULTS

Overall change in media trust

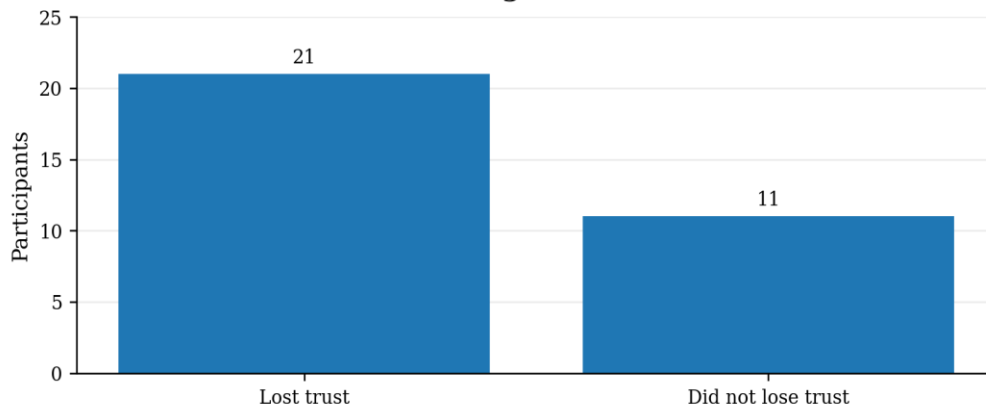


Figure 1: Overall number of participants who reported losing versus not losing trust in the media.

Media trust outcome by age group

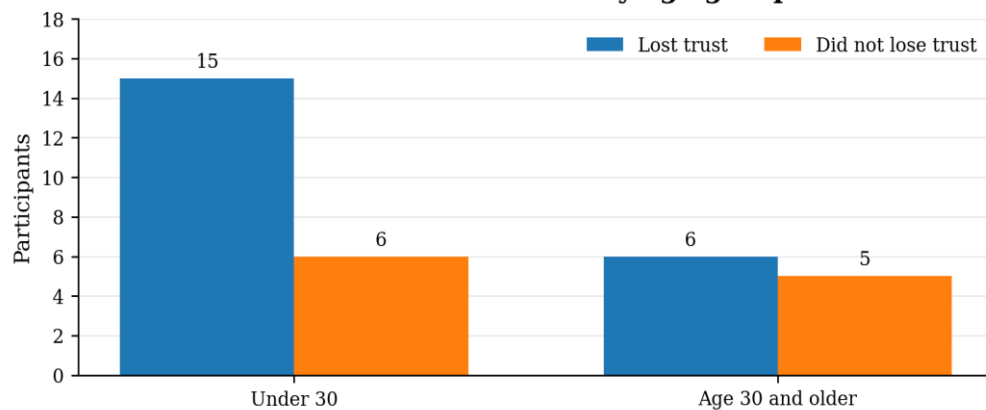


Figure 2: Media-trust outcomes among participants under age 30 and participants age 30 and older.

Media trust outcome by identification accuracy

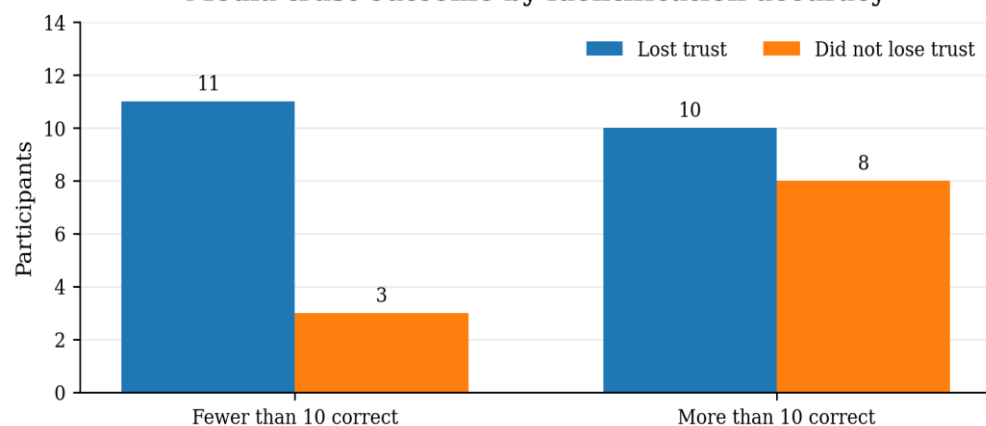


Figure 3: Media-trust outcomes by performance on the AI-generated-text identification tasks.

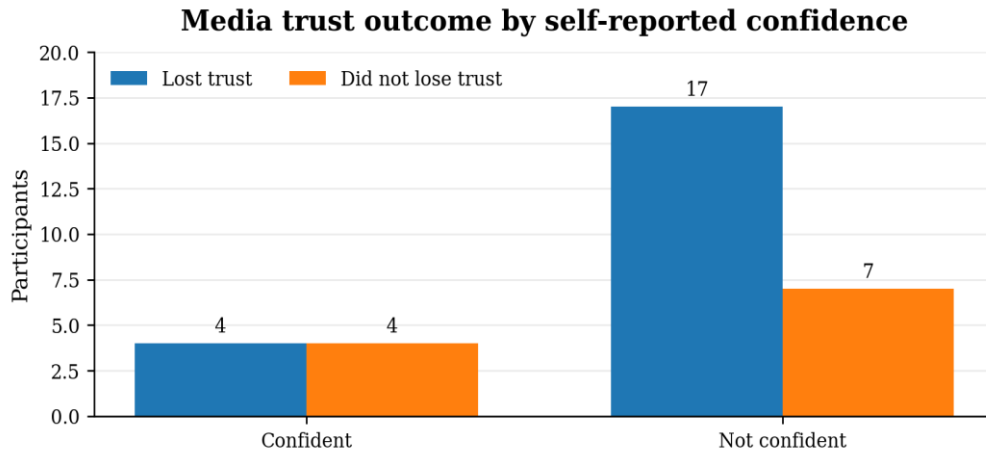


Figure 4: Media-trust outcomes by participants' self-reported confidence in identifying AI-generated text.

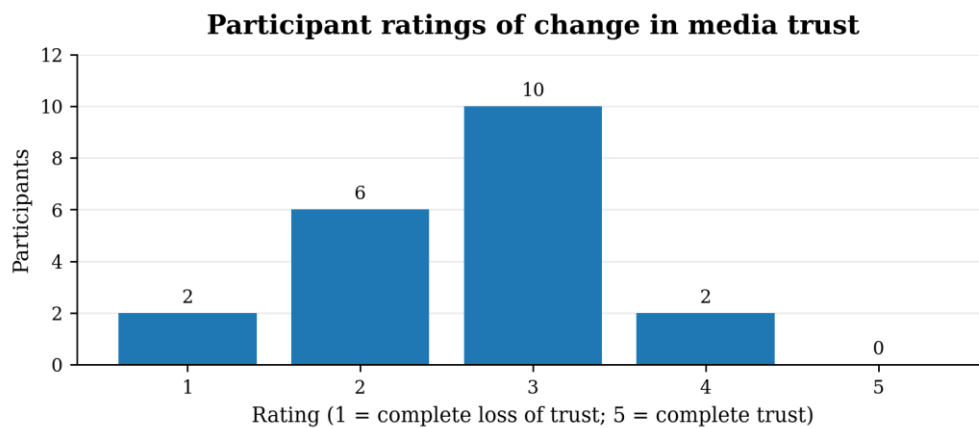


Figure 5: Participants ranked how their trust in the media evolved after the study on a scale from 1 to 5, with 1 indicating a complete loss of trust and 5 indicating complete trust.

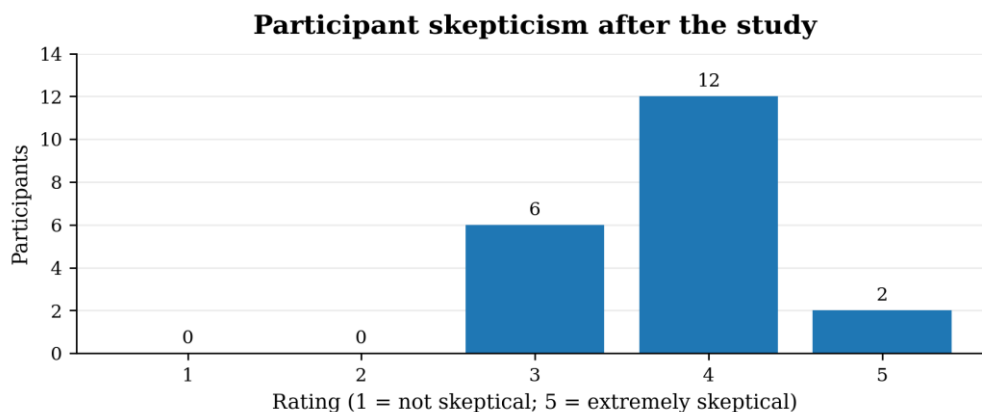


Figure 6: Participants ranked their skepticism toward online information after completing the study on a scale from 1 to 5, with 1 indicating no skepticism and 5 indicating extreme skepticism.

A total of 32 participants finished the pertinent survey item. Responses were classified dichotomously, with participants indicating lower trust awarded a value of 1 and

those indicating no reduction assigned a value of 0. The recorded percentage of participants indicating diminished trust was 0.656 (21 out of 32). A one-sample percentage test was performed to evaluate whether exposure to AI-generated texts reduced participants' faith in the media. The investigation assessed whether the percentage of individuals reporting diminished trust deviated significantly from 50 percent. Under the null hypothesis that the true proportion is 0.50, the test yielded a z-statistic of 1.41 and a two-tailed p-value of 0.159. The null hypothesis cannot be rejected because the p-value exceeds the standard alpha level of 0.05. The observed fraction is not statistically distinct from the expected value under chance. While most participants indicated diminished trust after the task, the effect did not achieve statistical significance in our group. This indicates that although the activity may affect certain individuals' perceptions of media trustworthiness, the existing data do not provide robust statistical evidence that exposure to AI-generated texts uniformly diminishes participants' confidence. You must regard these results with caution. The absence of statistical significance may be partially due to the limited sample size, which constrains statistical power. Subsequent investigations employing bigger and more heterogeneous groups, along with direct pre- and post-numeric assessments of media trust, would yield a more rigorous evaluation of the concept.

Additionally, the proportion of individuals under 30 who lost trust in the media was higher (0.714) than that of those over 30 (0.545). Additionally, all individuals who infrequently engage with the media saw a decline in trust in the media. However, among those who frequently engaged with the media, the proportion who lost trust was smaller (0.593). Participants who exhibited lower confidence in their ability to distinguish AI-generated texts from human-written texts before being administered the survey showed a greater proportion of diminished trust in the media (0.708) than those who were more confident (0.5). It was also found that participants with greater proficiency in distinguishing AI-generated from human-written texts had a lower proportion who lost trust (0.556) than those with lower proficiency (0.786).

CONCLUSION

There was not enough evidence to conclude that exposure to AI-generated texts and having to discern between AI-generated texts and human-written texts significantly decreases people's trust in the media. Additionally, future studies should use a larger, more diverse sample. With a bigger sample, the data will be more representative of the population. Even though this study didn't find a significant loss of trust after the study, when people engage with the media, they may want to consider that an article may have been generated by AI and take a moment to ensure the information is factual.

The proportion of individuals under 30 who lost trust in the media was higher than that of those over 30. This result is the opposite of what was hypothesized. Before data collection, it was thought that those under 30 were more likely to lose trust in the media than those over 30, since I believed that the younger generation is exposed to AI-generated texts at a higher frequency than the older generation. I thought that the more experience the younger generation has in identifying AI-generated texts, the more confident they would become and the better they would be at determining whether a text is AI-generated, which would lead them to lose less trust in the media. However, the results I found contradicted my belief: those under 30 lost trust in the media more frequently than those over 30. I

believe this finding is due to older people having more experience with the media in general than younger people. Even without the possibility of AI generating the texts they are reading, misinformation was still prevalent in the media, and they still had to determine if the text was real or fake. The more experience they have in detecting false information in the media, the more help the older generation needs with the new issue of determining whether the text was generated by AI or written by a human. They will have more confidence in their ability to identify misinformation since they can draw on past experiences engaging with the media and determining whether information was false or accurate. Additionally, many in the older generation reported being already distrustful of the media before the study. Participants over the age of 30 responded to the question, "Did this activity make you not trust the media in the future?" with, "I always question the media since there are many instances where a statement is made in the media headlines...but later corrected as a footnote," or, "I think I am already skeptical about the things that I read online, and this activity validates that skepticism." This activity wouldn't lead them to lose their trust in the media, as they already distrust it based on past experiences. Exposure to AI-generated texts would only confirm their preexisting distrust of the media. They were already aware of the false information in the media before this activity. Additionally, many of the participants over 30 years old didn't just indicate that they were already distrustful of the media, but also that they take steps to verify the source's credibility. Some participants responded to the same question as before with "I am aware that some things on the internet are not true, which is why I get my news from trusted media sources," and "I never trust anything unless it is in the major news outlets like The New York Times and CNN." Due to their distrust of the media, they have found sources they trust within it. They recognize the risks of AI-generated texts and misinformation in the media and take steps to ensure the media they read is factual by verifying that sources are trustworthy and credible. People in the younger generation didn't report a pre-existing distrust of the media, unlike older people. They don't have as much media experience to cause as much distrust. So when participating in this activity, they are less aware of false information in the media, which may have made them more aware of misinformation and AI-generated texts.

As I hypothesized, the proportion of those who frequently engaged with the media lost their trust less than that of people who infrequently engaged with the media. I believe this conclusion is true because people who engage more frequently with the media are more experienced at detecting AI-generated text and identifying false information. The more experience they have, the more likely they are to trust the media, since they can draw on past experience detecting false information and AI-generated texts when engaging with it. People who infrequently engage with the media don't have as much experience determining what information is false and what is AI-generated when they do engage.

Moreover, a greater proportion of participants who reported lower confidence in their ability to distinguish AI-generated from human-written texts before the survey was administered indicated diminished trust in the media, compared with those who were more confident in that ability. I believe that people with higher confidence are more likely to rely on their abilities to determine whether a text is AI-generated. So they are less likely to lose trust in the media because they believe they can identify AI-generated texts and use that information to decide whether to trust the source. People with lower confidence believe they cannot rely on their abilities to determine whether a text is generated by AI. Such doubts may make them more likely to lose trust in the media, since they believe that if they

come across an AI-generated article, they would be unable to identify it as such, and without this knowledge, they are more likely to believe the information is factual. So, since they can't rely on their abilities to prevent them from believing in potentially false information, they lose trust in the media.

It was also found that participants who were more proficient at distinguishing AI-generated from human-written texts had a lower proportion of those who lost their trust than those who were less skilled. I believe this is for similar reasons to why I believe perceived ability affected people's trust in the media, since those who are more proficient at correctly identifying AI-generated texts will most likely have higher confidence in doing so. Since AI-generated content is so prevalent in today's society, most people are likely to know whether they can correctly identify it as such and, if not, why. Even if they aren't consciously aware of whether they are proficient at this skill, they know that in the past they tended to be able to identify texts correctly or incorrectly. So I believe that those who are more accurate at identifying AI-generated content know their accuracy or that they have been accurate in the past, and they may believe that if they come across an article generated by AI, they will be able to correctly identify the text as such. Those who aren't proficient at identifying AI-generated content also know their inaccuracy or that they have been inaccurate in the past, so they may believe that if they come across an AI-generated text, they wouldn't be able to correctly identify it as such.

IMPLICATIONS

As AI progresses, it can produce increasingly lifelike images, movies, and texts that are more likely to mislead people. Although media consumers may be able to currently distinguish between AI-generated and human-created content, advancements in AI may make this distinction increasingly difficult in the future. Such advancements may lead more individuals to accept misinformation disseminated by AI-generated media. In particular, the younger demographic that engages most frequently with AI-generated material is at risk. The younger generation will be the ones making democratic decisions, working in the workforce, and raising the next generation of children. If they become more susceptible to believing false information, such behavior can be detrimental to society in the future. For instance, misinformation about the healthcare system can affect people's medical decisions. They may be misinformed about a treatment plan or medication, preventing them from receiving it when it could benefit their health. Additionally, the media is where people often will get their information about certain candidates for elected government positions. If someone spreads false information about a candidate, it may affect a voter's decision about whom to vote for. Such actions can lead to the election of someone who isn't the best-qualified for a position that affects the lives of many people. This study addresses an issue that is likely to worsen, underscoring the need for further research to understand the implications of AI. A lack of faith in the media may result in significant societal consequences. If individuals lose faith in the media, they may stop consuming it or become skeptical of the information. If they stop consuming media because they cannot trust the information, it may be difficult for them to stay up to date on current events and to learn about critical aspects of their lives, such as society, politics, economics, and healthcare. In the future, research with a much larger and more diverse sample could provide statistically significant evidence that exposure to AI leads to a loss of trust in the media. Also, experimenting with other forms of AI-generated content, such as videos, audio clips, and images, may help us better

understand how AI-generated content affects people's trust in the media. Videos and audio may be more emotionally persuasive and even harder to verify than texts. They may lead to more distrust in the media if more people are deceived by videos and audio.

REFERENCES

- Ahmed, S., & Tan, H. W. (2022). Personality and perspicacity: Role of personality traits and cognitive ability in political misinformation discernment and sharing behavior. *Personality and Individual Differences*, 196, 111747. <https://doi.org/10.1016/j.paid.2022.111747>
- Atske, S. (2025, November 20). Americans' Social Media Use 2025. Pew Research Center. <https://www.pewresearch.org/internet/2025/11/20/americans-social-media-use-2025/>
- Babajide Osatuyi, & Dennis, A. R. (2024). The strength of weak ties and fake news believability. *Decision Support Systems*, 184, 114275-114275. <https://doi.org/10.1016/j.dss.2024.114275>
- Baig, E. C. (2022, August 17). 10 Ways to Spot Fake Videos and Falsehoods on the Internet. AARP. <https://www.aarp.org/personal-technology/identifying-disinformation/>
- Calabrese, E. (2020, January 7). 5 ways to spot disinformation on your social media feeds. ABC News; ABC News. <https://abcnews.go.com/US/ways-spot-disinformation-social-media-feeds/story?id=67784438>
- Chaudhuri, N., Gupta, G., & Popovič, A. (2025). Do you believe it? Examining user engagement with fake news on social media platforms. *Technological Forecasting and Social Change*, 212, 123950. <https://doi.org/10.1016/j.techfore.2024.123950>
- Goyanes, M. (2026, March). Trust in AI news, AI literacy, and the mediating role of artificial intelligence attitudes: A longitudinal study across diverse societies. *Science Direct; science direct*. <https://www.sciencedirect.com/science/article/pii/S2949882126000307>
- Maes, P. (2025). People Overtrust AI-Generated Medical Advice despite Low Accuracy - MIT Media Lab. MIT Media Lab. <https://www.media.mit.edu/publications/NEJM-AI-people-overtrust-ai-generated-medical-advice-despite-low-accuracy/>
- Makowski, D., Te, A. S., Neves, A., Kirk, S., Liang, N. Z., Mavros, P., & Chen, S. H. A. (2025). Too beautiful to be fake: Attractive faces are less likely to be judged as artificially generated. *Acta Psychologica*, 252, 104670. <https://doi.org/10.1016/j.actpsy.2024.104670>
- Melchior, C., Warin, T., & Oliveira, M. (2025). An investigation of the COVID-19-related fake news sharing on Facebook using a mixed methods approach. *Technological Forecasting and Social Change*, 213, 123969. <https://doi.org/10.1016/j.techfore.2024.123969>
- Rajesh, D., & Vijayashri Gurme. (2025). Artificial Intelligence and Consumer Behaviour on Social Media: A Study of Personalization, Trust, and Privacy. *Measurement Digitalization*, 100021-100021. <https://doi.org/10.1016/j.meadig.2025.100021>
- Traberg, C. S., & van der Linden, S. (2022). Birds of a feather are persuaded together: Perceived source credibility mediates the effect of political bias on misinformation susceptibility. *Personality and Individual Differences*, 185, 111269. <https://doi.org/10.1016/j.paid.2021.111269>
- Trandabăț, D., & Gifu, D. (2023). Discriminating AI-generated Fake News. *Procedia Computer Science*, 225, 3822-3831. <https://doi.org/10.1016/j.procs.2023.10.378>