

Open
AccessCheck for
updates

REVIEW ARTICLE

Section: *Culture, Media & Film*

Ethical challenges of using artificial intelligence in news and television program production

Alaa Makki^{1*}, Fawzia AlAli¹, Omer Jawad Abduljabbar¹ & Majed Z Hatem Alkindi²¹Department of Mass Communication, College of Communication, University of Sharjah, United Arab Emirates²Department of Mass Communication, College of Arts, Education, and Social Sciences, Abu Dhabi University, United Arab Emirates* Correspondence: r.alsadi@psau.edu.sa

ABSTRACT

Machine learning is now used in every part of news and TV shows. Most of the time, when people discuss if automation is good or bad, they are not talking about a specific task. This piece discusses the moral issues that arise when AI is used to create TV shows and news. This is why it examines how AI is used to create shows, if viewers can still understand their purpose, and if editors can still be held responsible. The way AI is used in production and communication is what creates ethical risk, not the technology itself. Thoughts on gatekeeping, media duty, and sociotechnical systems support this claim. An online poll experiment with 463 people watched short news stories and TV shows that were different based on how AI was used in the production process. This was done to test the claim. These jobs included material made by humans, speakers made by AI, visual rebuilding made by AI, the amount of human editing control, and the reveal condition. People who took part said how much they believed in the situation, how likely they were to be manipulated, how willing they were to watch and share, and how responsible they thought the situation was. The most trusted content was that which was made by humans with little or no help from AI. When AI's part became clearer, especially in AI presenters and AI-made visual reconstructions, people lost a lot of faith in the content. People didn't always trust each other more when they told the truth. Trust and responsibility went up the most when disclosure was paired with named human editing control, not when disclosure was shown on its own. These results back up a model of AI ethics in news that is based on how things are done. The level of risk relies on where AI is used in production and whether the employee is clearly responsible for their actions. AI isn't needed at all. Numerous people have talked about what this means for running newsrooms, creating rules for AI exposure, and teaching journalism.

KEYWORDS: Artificial intelligence, journalism ethics, television production, generative AI, newsroom automation, synthetic media,

audience trust

Research Journal in Advanced Humanities

Volume 7, Issue 3, 2026

ISSN: 2708-5945 (Print)

ISSN: 2708-5953 (Online)

ARTICLE HISTORY

Submitted: 07 June 2026

Accepted: 25 August 2026

Published: 04 September, 2026

HOW TO CITE

Makki, A., AlAli, F., Abduljabbar, O. J., & Alkindi, M. Z. H. (2026). Ethical challenges of using artificial intelligence in news and television program production. *Research Journal in Advanced Humanities*, 7(3). <https://doi.org/10.58256/z5re3777>



Published in Nairobi, Kenya by Royallite Global, an imprint of Royallite Publishers Limited

© 2026 The Author(s). This is an open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

1. Introduction

News is now found, checked, written, drawn, edited, sent, and judged by AI. Nobody just uses it to write text or papers, though. AI can now help newsrooms collect information, translate it, summarize it, make new information, edit videos, repackage them for social media, research audiences, and make suggestions. They are morally distinct since they involve working on the movie at various times and with various levels of risk, exposure, and editing impact. The tool a producer uses to record is not as illegal as the visual reconstruction made by AI that is added to a news show. A fake speaker who talks like a writer is also not as trustworthy as a language helper. Some individuals don't grasp how editing, technical, visual, and organizational methods all function together when AI is seen as a solitary device (Broussard et al., 2019; Lewis et al., 2019; Wu et al., 2019). There are ethical concerns about not only whether AI is used or not, but also where it is used, how well it can connect with other things, and whether people can still be held accountable. Machine learning is now used in all areas of news and TV shows. Often, when people talk about whether automation is good or bad, they are not referring to a specific task. This piece talks about the moral problems that come up when AI is used to make TV shows and news. This is why it looks at how AI is used to make shows, if viewers can still understand their purpose, and if editors can still be held responsible. When people read something, they might want to know who wrote it or if the name of the author is clear. Watching TV, some people may question whether a person, voice, scene, or video recording is real, fake, rebuilt, or agreed to. Because of these changes, people are more worried about picture rights, deepfakes, visual edits, and the loss of broadcast standards that are based on facts (Ndlovu, 2024; Raemy et al., 2025; Thurman et al., 2021). Some people might think something is true before they check it or change it if the news makes it look authentic. Most likely, this will happen when there is news about the government, a war, public health, crime, or a disaster.

AI makes people less sure that writers are trustworthy, can be checked, and only want what's best for the public. What people believe about something's reliability relies on who wrote it, the type of writing, the amount of human input seen, and where the study took place (Liu & Wei, 2019; Wang & Huang, 2024; Wölker & Powell, 2021; Zheng et al., 2018). Recent research shows that being open is not an easy answer. While viewers might like transparency, it could reduce trust if it seems like journalists are relying too much on computers without enough human oversight (Jia et al., 2024; Toff & Simon, 2025; Valenzuela et al., 2026). A lot of people might not have used AI to write, record, make scenes or faces or change how someone shared something if it wasn't clear how it helped. That's the same thing you should think about when you use AI: what people missed, who gave permission for the release, and how to fix mistakes.

A number of studies, including Carlson (2018), Cools & Koliska (2024), Helberger et al. (2022), and Porlezza (2023), have looked at automation, rules, the use of AI in newsrooms, and how to handle AI. People still talk a lot about easy ideas like being fair, honest, responsible, and open when they talk about AI ethics. These ideas should not only be thought about at the company level, but they should also be put into action. There are various interpretations of the word "accuracy" in a report created by AI, a fake recreation, and a customizable spread. A process method can help you understand why the same moral idea is always right.

The study looks at the moral problems that arise when AI is used to make news and TV shows. It is looked at in terms of how edits, fake news, being open, and public trust are all connected. Some people think the worst things about AI come from the way it is shown on TV and in the news. This has been called into question. How do people's ideas about who should edit change when they learn that AI was used? People may or may not trust AI or something written in a certain way if someone is in charge of it. The second idea says that people are more likely to trust AI when it reads or records in the background. When AI is used, like with fake pictures or AI programs, it makes people less likely to believe it. People are more likely to believe it if someone checked it and it is clear who did it than if no one is named but it is clear who did it.

2. Literature Review and Theoretical Positioning

2.1 AI in news and television production workflows

AI isn't just a tool for writing. People can also get news from it. A TV station or newsroom does many things, like find stories, check sources, and help with data analysis. They can help you read, write, translate, and sum up when you are studying. When you review, these can help you write scripts, leads, copy, or different types of stories for different platforms. In the TV business, AI is also used to make images, voices, reconstruct images,

act as hosts, tag metadata, and make suggestions. All of these steps are important from an ethical point of view because they all have something to do with proof, authorship, and how the audience sees the work. There are back-end tools that can change the truth, and presentational tools that can change how the reader sees the writer's looks and trustworthiness. We all know that writers don't always use AI the same way for every job. They use it in a skilled way that takes into account what the group can do and how dangerous they think it might be. This means that the government has less power to run things.

2.2 Ethical challenge domains

When you compare how things are made in the real world to how AI helps make things, you need to think about some moral issues. People might not get the whole picture or understand it correctly if they use AI-made recaps, translations, or recommended sources without checking the original source or making sure they know what it means. It is not fair or right to show certain groups or back up assumptions when these things happen with training data, automatic ideas, or visual generation. But this can also happen when these things make regular traditional views seem strange. It's not always clear what "AI use" means. What does it mean? It could mean making up songs and movies or writing down your own thoughts. It's hard to be open and honest because a general disclosure might not reveal much about the writing help. A lot of people work with editors, like writers, directors, software providers, and platform systems, as Carlson (2018) and Lewis, Sanders, and Caremody (2019) point out. This makes it hard for them to be responsible.

It is possible to make fake voices, faces, and sounds, as well as to rebuild them and use them in new ways. That's why it's important to be able to watch TV alone and with permission. A video that looks like it was made to teach can be used as proof, even though it wasn't made to teach. When you mix real and fake news, this new problem shows up. Losing your job and professional identity is linked to how the media changes ideas about what is right and wrong, skilled, and powerful (Milosavljević & Vobič, 2019). In the end, fake or automatic content that changes what people know about politics, how they talk about situations, or how much they believe what everyone says is bad for democracy and the public good. There doesn't seem to be any difference between these places on this list. In the event that a fake reconstruction is used in a crime story without being talked about, it may not be true, get permission, be checked, or gain public trust. An automatic piece could also be false, have unfair edits, or be harmful to society (Raemy et al., 2021). You shouldn't just choose a production method based on its name, the press's history, or what the seller says.

2.3 Audience trust, disclosure, and human oversight

How much people trust the media depends on what they think about AI and whether they think it is helping people or replacing them. What kind of bias people think it has, who wrote it, and where the information came from all affect how people feel about it (Liu & Wei, 2019; Wang & Huang, 2024; Wölker & Powell, 2021). AI could be used in the back end if it seems to make things run more smoothly without taking the place of human judgment. People are less likely to believe things when they hear AI voiceovers, see fake pictures, and listen to fake speakers instead of real people, proof of skill, and human presence. Being honest is fine, but only if it makes sense. Even if there is a simple sign, people might not understand what is set to automatically, what has been checked, or who is still in charge. Some people say that too much exposure can make things less reliable when we talk about editing distance or machine freedom (2021). The government can't just say it will look out for people. There are rules to follow, sources to check, permissions for images and sounds, and names of people who need to fix things. Let people know you can change things and are open. This will make them trust you more. This is especially true when people are asked to trust video footage as proof instead of just reading hard-to-read news stories about public problems.

2.4 Theoretical positioning

This research uses the gatekeeping theory, the sociotechnical systems theory, and the media responsibility theory. It's smart to put up gates. This means that facts are checked by someone or something before they are shared. These tools can change the sources that are used and how they are grouped, changed, organized, and shared. It will be open to everyone since people and tools will be able to work together (Wu et al., 2019). When it comes to media accountability theory, there are four main ideas. These are duty, control, power, and work-related power.

These days it's harder to use AI because systems can make mistakes that people who watch or write can't see. Lawyers and editors at the media companies that made these mistakes are still to blame. The sociotechnical systems theory says that people are morally endangered by more than just technology. People, platforms, tools, providers, users, habits, and rules all work together to make it happen. This article agrees with this point of view because the same AI tool can raise different moral issues depending on how it is used, how public it is, how it is managed, and how it is changed (Helberger et al., 2022; Jones et al., 2022). Then there's the connection between relationships and risk.

2.5 Research questions and hypotheses

Findings show that just asking if a newspaper uses AI is not enough to figure out what the social effects of AI in journalism are. The important parts of the manufacturing process that use AI, how important that role is, and how people see it should all be looked at. You ask Respondent 1 what kinds of AI applications they think pose the biggest threat to society. In light of the fact that AI is involved in creating material, the second question asks how people think writers should manage their duties. What kind of openness, human control, and material style do people have faith in you? That's the third question. Depending on how the audience sees things and the idea of accountability, research shows that people are more likely to trust back-end AI applications like transcription or translation than front-end AI applications like fake voices or images. Having clear exposure and clear human editing control will make people more confident (H2).

3. Methodology

3.1 Research design

The goal of this study is to find out how making AI's role in making news and TV shows more clear affects people's trust, belief, and willingness to take responsibility for their actions. People were given a short written story about a news or TV show that was made with a mix of human editing control, conditions for disclosure, and production jobs for AI. They had to write about how that made them feel afterward. The tests used different AI output jobs, notification conditions, and levels of control that were based on current talks about how to regulate AI in newsrooms and media and the academic literature we looked at above (Helberger et al., 2022; Porlezza, 2023) That is, the changes are based on differences that are already being talked about in business and government, not on categories that were made up just for the study. This method lets the study check, at the level of how people see it, whether ethical risk changes depending on the work being done, as the literature review shows. It also leaves open the important route for future study to look directly at choices made by journalists and policy papers.

3.2 Participants and procedure

News readers were asked to join an online group. Quota sampling was used to find people in the group who were about the same age, gender, and level of education as people who read news in general. The last sample was made up of responses from 463 people. The first look at the data showed that none of the key measures for the study were missing. After providing informed consent, participants reported demographic information, news interest, familiarity with AI, and political interest as control variables, and were then randomly assigned to one of the experimental conditions described in Table 1.

Table 1. Experimental Conditions and Measures

Factor	Levels	Description
AI production role	6 levels	Human-made content; AI-assisted script with human editor; AI-generated script with human review; AI-translated or AI-voiced segment; AI-generated visual reconstruction; AI presenter/synthetic anchor
Disclosure condition	4 levels	No disclosure; minimal AI-use label; detailed AI-use disclosure; detailed disclosure naming a responsible human editor
Human editorial oversight	3 levels	No described oversight; general oversight; named, editorial-level oversight

Outcome measures	11 measures, 1–7 scales	Audience trust; perceived credibility; perceived authenticity; manipulation concern; willingness to watch; willingness to share; perceived professionalism; disclosure clarity; perceived accountability; bias concern; privacy concern
------------------	-------------------------	---

3.3 Measures and experimental manipulation

There were six different ways that the AI production role could be changed: content made by humans, scripts written by AI with human review, scripts written by AI with human review, segments translated or voiced by AI, visual reconstructions made by AI, and AI presenters or synthetic anchors. Four levels of disclosure were changed: no disclosure, a minimal AI-use label, a detailed AI-use disclosure statement, and a detailed disclosure statement that named a human editor who was responsible for the work. Three levels of human editing oversight were used to code them: no specified oversight, general oversight, and named editorial-level oversight. These changes match the risk-tiered categories used in current discussions about AI policy in newsrooms and broadcasters (Helberger et al., 2022; Porlezza, 2023) and the theoretical difference between back-end production support and AI use that is aimed at the audience and is sensitive to their opinions.

The participants rated on a scale from 1 to 7 how much they trusted the audience, how credible they thought the scenario was, how real they thought it was, how willing they were to watch and share, how professional they thought it was, how clear they thought the disclosures were, how accountable they thought the participants were, how concerned they were about bias and how concerned they were about privacy. It's shown in Table 2 how the conceptual coding reasoning linked different production-stage AI to the moral problems they have to solve.

Table 2. Conceptual Coding Scheme Linking Production Stages to Ethical Risk Domains

Production stage	Representative AI use	Associated ethical risk	Risk indicator	Illustrative application
Research	Source discovery	Accuracy / source bias	AI-cited source unverifiable	Source verified vs. unverifiable
Scripting	Automated rewriting	Editorial distortion	Meaning changed from original	Distortion rated low/medium/high
Visual production	Generated image/video	Authenticity	Realistic scene without real footage	Disclosed vs. not disclosed
Voice/presenter	Synthetic voice/AI anchor	Identity/consent	Human likeness simulated	Consent present vs. absent
Distribution	Personalization	Fairness/public interest	Unequal exposure to news	Risk level (low/medium/high)

3.4 Analytic strategy

What was important were how AI was used in production, how disclosure worked, and how much editing was done by humans. Content style and topic sensitivity were two other coded factors. Gender, age, level of education, place of residence, interest in news, knowledge of AI, and interest in politics were all used as control variables. The outcome factors were audience trust, perceived sincerity, perceived trustworthiness, concern about influence, readiness to watch, willingness to share, perceived professionalism, disclosure clarity, perceived responsibility, concern about bias, and concern about privacy. The scale ranged from 1 to 7. We used ANOVA to see how trust levels changed for various AI jobs, and hierarchical multiple regression to test H1 and H2. The first thing that was done was to look at the role and state of AI creation. The second step was to look at the fact that there would be human editing control. The third and last step was to look at how the AI role disclosure and AI role oversight terms work together.

3.5 Validity, reliability, and research ethics

Initial tests were done on the experimental data to make sure that subjects could tell the difference between situations where things were made by humans, with help from AI, and when they were made by machines. Respondents gave their informed consent, and their answers were collected anonymously. Artificial intelligence (AI)-related stimuli were clearly presented as part of a research exercise rather than as real news, and no fake content was shared outside the study.

4. Results

4.1 Sample characteristics and condition balance

The final sample was made up of 463 respondents, and there were no missing values for any of the important study variables. Human-made content ($n = 73$), AI-assisted script with human editor ($n = 77$), AI-generated script with human review ($n = 79$), AI-translated or AI-voiced segment ($n = 77$), AI-generated visual reconstruction ($n = 77$), and AI presenter/synthetic anchor ($n = 80$) were all used in equal numbers in the experiments. No disclosure ($n = 111$), minimal AI label ($n = 118$), detailed AI-use disclosure ($n = 118$), and detailed disclosure with a named human reviewer ($n = 116$) were all approximately equal chances of occurring. With this mix, it's simple to see how AI jobs that help with the back end are different from those that help with the front end and talk to users.

4.2 Audience trust across AI production roles

These results show that perceived ethical risk, which was measured by audience trust, was not spread out evenly across AI roles. Based on the mean trust score, human-made content got the highest score (5.48), followed by AI-assisted script with human editor (5.22), AI-generated script with human review (4.95), AI-translated or AI-voiced segment (4.78), AI-generated visual reconstruction (4.43), and AI presenter/synthetic anchor (4.31). At the public impression level, this trend supports H1: uses of AI that were easy to see got less trust than uses of AI that were harder to see in production jobs. Findings from a one-way ANOVA showed that trust levels were significantly different between AI job conditions ($F = 38.39$, $p < .001$). These results are shown in Figure 1.

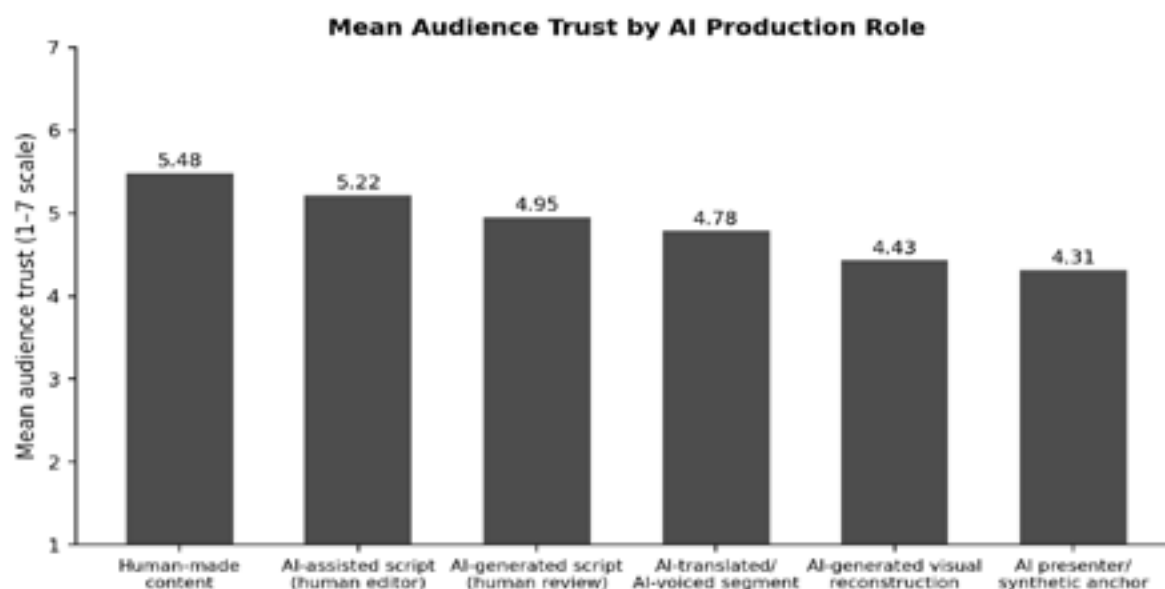


Figure 1. Mean Audience Trust by AI Production Role

The same pattern can be seen in worries about authority and power. The most real content was created by people ($M = 5.60$) and a script with AI help that was edited by a person ($M = 5.10$). The least real content was made by AI and was edited by a person ($M = 3.96$). A lot of people were more worried about content made by AI (4.43 million) than content made by humans (2.66 million). The things that worried people the most were visual restoration done by AI and AI presenters and fake anchors. These findings show that watching risk is mostly linked to how people see reality, how credible fake news is, and how much fake news there is.

4.3 Perceived accountability across production and governance conditions

How does the public feel about editorial responsibility when AI plays a big role? That's what RQ2 is all about. The perceived-accountability measure is a simple way to test this question. The most reliable content was made by a person ($M = 6.45$), followed by a script written by an AI with help from a human editor ($M = 5.65$), a script written by an AI with help from a human reviewer ($M = 5.22$), a segment translated or voiced by AI ($M = 4.97$), a visual reconstruction made by AI ($M = 4.03$), and an AI host or anchor ($M = 4.01$). People are much less likely to hold AI responsible when it can be seen in things like recognizing presenters or reconstructing images.

Watching and telling others about things changed how responsible people thought others were as well. It took the least amount of responsibility to not disclose ($M = 4.53$), and it took even less to disclose ($M = 4.67$). It made people more responsible ($M = 5.00$) to know more about how AI was used, and it made them even more responsible ($M = 5.94$) to know more about how AI was used with a named human boss. Oversight at all levels ($M = 3.7$), specific oversight ($M = 4.91$), and specific oversight with a named person in charge ($M = 6.09$), as well as overall oversight, all raised the level of responsibility. These results show that giving isn't socially good unless it's clear that you have to.



Figure 2. Mean Perceived Accountability by AI Production Role, Disclosure Condition, and Human Oversight Level

4.4 Regression analysis of trust, disclosure, and oversight

How much trust people had in the AI depended on its role, the type of information it shared, and who was watching over it. It was easiest to see the trust cost when AI tasks were clear. There was an average trust score of 4.37 for high-visibility AI situations, such as AI-generated visual reconstruction and AI presenter/synthetic host. This was less than the scores of 4.99 for AI conditions that were less obvious and more than the scores of 5.48 for content made by humans. The Welch tests showed that people rated high-visibility AI jobs much lower than less visible AI roles ($t = -9.19$, $p < .001$) and lower than content made by humans ($t = -12.99$, $p < .001$).

It was harder to say what effect disclosure had. A confidence score of 4.82 meant that there was no disclosure, a score of 4.64 meant that there was limited AI labeling, a score of 4.85 meant that there was detailed disclosure of AI use, and a score of 5.10 meant that there was detailed disclosure with a named human reviewer. This pattern supports the idea that disclosure alone is not enough. If the term doesn't specify what AI accomplished or who is still in charge, it can increase people's fear. Being transparent was insufficient to foster confidence; "human control" was also required. The results of the hierarchical regression are shown in Table 3. These findings partly back up H1 and H2. There was a big drop in trust for AI-generated visual reconstruction and AI presenter/synthetic anchor compared to human-made content in the first step ($B = -0.874$, $SE = 0.148$, $p < .001$), but not for AI-assisted script with human editor ($B = -0.156$, $SE = 0.120$, $p = .194$). Little information was found to be bad ($B = -0.199$, $SE = 0.081$, $p = .015$). For the first step, having a person edit the work was good, but not very important ($B = 0.146$, $SE = 0.088$, $p = .096$). In the second step, it became significant ($B = 0.213$, $SE = 0.098$, $p = .031$). In the last and third step, weaker interaction terms were added. It wasn't important for the AI role to reveal or oversee anything. The final model in Table 3 backs up the main-effects version of the argument more strongly than the moderation version. Its coefficients can be seen in Figure 3.

Table 3. Hierarchical Regression Predicting Audience Trust (Final Model)

Predictor	B	SE	p
AI-assisted script (vs. human-made)	-0.164	0.131	.213
AI-generated visual reconstruction (vs. human-made)	-0.918	0.224	< .001
AI presenter/synthetic anchor (vs. human-made)	-1.044	0.223	< .001
Minimal disclosure (vs. no disclosure)	-0.199	0.081	.015
Detailed disclosure (vs. no disclosure)	-0.140	0.099	.161
Human editorial oversight	0.137	0.113	.228
AI role × disclosure condition	-0.557	0.288	.054

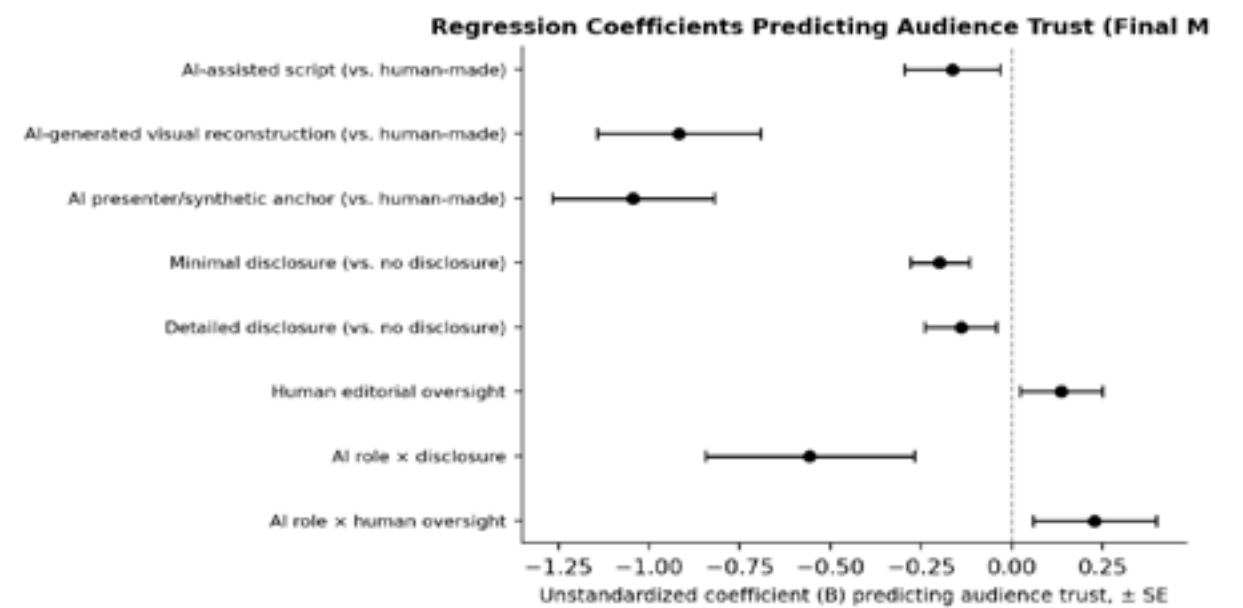


Figure 3. Regression Coefficients Predicting Audience Trust (Final Model)

4.5 Integrated findings

All of the research have one thing in common: people think there is increasing moral risk as AI grows more real and less responsible. AI seems to modify or edit; it creates realistic video content using fake noises, graphics, or presenter identities; disclosure is not accompanied by a clearly recognized human reviewer. This is obviously dangerous. This study found that people lost the most trust when AI was used as a host, anchor, or to recreate images. On the other hand, there was the greatest degree of trust and responsibility when anybody could join and people were given power.

5. Discussion

5.1 Main contribution: from general AI ethics to workflow-based media ethics

Its primary contribution to ethics research is the shift in emphasis from generic AI regulations to workflow-based media ethics. Users did not see AI in TV and news as a single, uninteresting category, according to a survey. When AI was used as a helper, such as when individuals modified scripts that AI had assisted in writing, people still put a lot of confidence in it. When AI voices or phony images were employed to demonstrate AI, this figure fell precipitously. Further evidence demonstrates that the ethical danger varies according on how the tool is utilized, how effectively it communicates, and how much AI seems to dispute the news. Automation has already altered how people consume media, make decisions at work, and communicate with machines (Carlson, 2018; Lewis et al., 2019; Wu et al.-2019). This research bolsters the notion that public perceptions of AI are socially significant, particularly in the context of television, where credibility is influenced by voice, image, presence, and realism.

They also demonstrate that excessive use and danger are not synonymous. People may sometimes question the veracity of the material due to routine back-end tasks like internal reports, transcripts, or translation. But this doesn't always happen. Not as often, but more trusting, are AI speakers, made sounds, and computer-generated rebuilding. This is because they look more like real-life proof. The less responsible AI is, the more it works directly with customers instead of helping with production. There should be more than just the name of the tool used to measure risk in AI at the production level, according to the study. It also checks out where the work is kept, how clear it is, what the editor thinks, and who is in charge of it.

5.2 The transparency problem

The results show that we need to have a better idea of what it means to be open. People don't always trust you when you're honest. People were less likely to trust AI naming when it was seen in small amounts than when it

was seen with full human control. A new study says AI can make people feel like they're stuck. People may want to know everything, but if they get it in a way that is hard to understand or is done for them, they might not use their brains as much (Toff & Simon, 2021). The guidelines require that short labels, such as "AI was used," be unambiguous. However, such guidelines do not specify if AI created an image, captured a voice, authored a screenplay, translated a remark, replicated an interview, or molded the final edit.

The main issue is that openness speaks a lot about using technology to participate, but it doesn't specify who is changing. A person would believe that the corporation trusted a flawed system if they just learned that AI was employed and not what was examined. If disclosure reveals what AI did, who reviewed it, and who approved its distribution, we may hold individuals accountable rather than just cautioning them. In a perfect world, showing people how to use tools would also make it clear who is in charge. Which means that sharing that is important is specific, based on risk, and has to do with keeping track of named people.

5.3 Why television production raises higher ethical stakes than text news

TV shows bring up more social issues because they do more than just show facts. Some of the things it gives are being there, truth, feeling, urgency, and public power. This is possible with the help of AI hosts, fake sounds, movies that have been rebuilt, and pictures that look like they came from an old file. It is possible for things that didn't happen to look like they did. People who are used to seeing pictures on TV and hearing the way a show talks as proof of reality may see audiovisual realism as proof as well. People rate self-made movies and AI sounds by how real, trustworthy, and involved they seem (Ndlovu, 2024; Thurman et al., 2021).

TV shows don't just have fake parts; they also break moral rules. These kinds of graphs and pictures have been used in documentaries for a long time. Whether or whether individuals are able to distinguish between simulations and actual evidence is the more crucial issue. Is it possible for the group to impersonate a voice, face, identity, or emotional performance? It is becoming more difficult to distinguish between false news and actual news due to deepfakes. According to Raemy et al. (2021), this makes it simpler to undermine democracy, mistreat individuals, and intrude into their personal life. AI must safeguard more than just the news. It must safeguard the host's power in addition to the images and audio.

5.4 Governance implications for newsrooms and broadcasters

This shows that AI shouldn't be used just one way. Instead, the government should be based on how much is being made. People who make things and newsrooms should first make a plan for AI that includes different amounts of risk. With little risk, it can be used for many low-risk tasks, such as writing, word check, information help, and internal process support. With a medium level of risk, you could translate, summarize, write headlines, or get help with a script. These can change the words used, the tone, and the meaning. Using AI to make fake sounds, faces, shows, scenes, and news stories that are written by computers is risky. We shouldn't allow fake eyewitness videos, secret fake rebuilding, the use of fake public figures without permission, or the use of fake material in areas like politics, war, crime, public health, elections, or disaster reporting. There are times when you shouldn't be able to do or let these things happen.

On top of that, the government should make everything binding. Editors should have to agree to stories that use AI, and directors should have to agree to changes to sounds or pictures. Additionally, using a false name, face, or voice is improper and should be investigated by the authorities. According to Helberger et al. (2022) and Porlezza (2023), general transparency is not as crucial as duty, penalty, and responsibility. This makes sense. The degree of hazard should determine how much information is disclosed. If work tools don't alter their meaning or how others perceive them, they may not need to be made public. If the report is short, AI may need to compose it. You need a lot of images, sounds, and extravagant presenters for lengthier disclosures that highlight both the function of AI and human assessment. Fourth, in order to address the shortcomings of AI, we must reconsider how issues are reported and resolved. False information, deceptive images, unlawful use of likenesses, and artificial intelligence-generated material should all be reportable. Some people can trust writers, sites, editors, or service providers because there are other people who keep an eye on them.

5.5 Implications for journalism education and professional training

AI skills should not only be taught as a way to learn how to use AI, but also as a way to teach ethics and edits.

Kids should learn to read and write as soon as possible. They should also know how to check the reliability of news sources, spot slanted information, write warnings, understand copyright, protect their privacy, and give permission for their voice and picture to be used. Also, people in both college and the workplace should get used to making quick choices, since AI often fails when speed is more important than proof. AI tools shouldn't become too important for journalists. Instead, they are taught how to keep an eye on, question, record, and explain the decisions that machines make that help make things.

5.6 Limitations and future research

Issues with this study exist since it only uses one type of survey experiment with fake news and TV situations. The purpose of the test cases was to show how ongoing discussions about how to run the press and communication channels affect different types of production. However, they fall short of capturing the complexity and intricacy of a functional production system. By examining AI policies in newsrooms and media corporations in more detail, future research must improve on this approach. It must have in-depth discussions with reporters and producers, observe newsrooms over time, compare newsrooms across cultural boundaries, and test concepts using more authentic audiovisual content. You may need a more engaging presentation to demonstrate the moral significance of AI in TV creation.

6. Conclusion

AI has been at the center of many ethical disputes in the press. There are moral dangers at every level of the procedure. Not whether AI should be utilized in news and television, but how it alters things like who is in charge, how much people trust, and who authored what each time the meaning changes. The worst AI-related incidents don't always occur. People often don't alter their perceptions of the news whether they record, proofread, receive translation assistance, or just summarize. However, if these applications aren't examined, errors or prejudice may result. Despite this, AI hosts, voice actors, created images, reconstructed scenes, and self-playing video tales may all alter television.

These discrepancies are supported by what the study revealed. Even though it was obvious that individuals had altered the scripts, many people trusted human-written screenplays more than AI-written ones. People lost faith much faster when AI was utilized in public media, particularly to host and reconstruct images. People dislike AI for more reasons than just that; it's also phony. Even worse for them is when AI seems to fill positions that people ought to have, such as being present, a witness, an investigator, or a qualified judge. People are more concerned about ethics when they don't know what's going on or when there isn't a clear person or editing unit recognized as accountable for the final result.

If newsrooms and the media wish to control AI, they need to do more than simply claim to be transparent and responsible. Things should be evaluated based on their effectiveness and level of hazard. It is necessary to ensure the effectiveness of tools for tasks that don't involve a lot of danger. The definition, operation, and applications of medium-risk change tools need to be discussed. For high-risk fake video tools, there should be a clear warning, authorization checks, editing approval, and a review by a person with extensive legal or moral understanding. Only under extreme circumstances and when there is a valid basis to do so such as during a crisis, election, war, criminal activity, or public health emergency should it be carried out. For example, you shouldn't utilize AI-generated presentations or robots that reassemble objects.

References

- Broussard, M., Diakopoulos, N., Guzman, A. L., Abebe, R., Dupagne, M., & Chuan, C.-H. (2019). Artificial intelligence and journalism. *Journalism & Mass Communication Quarterly*, 96(3), 673–695. <https://doi.org/10.1177/1077699019859901>
- Carlson, M. (2018). Automating judgment? Algorithmic judgment, news knowledge, and journalistic professionalism. *New Media & Society*, 20(5), 1755–1772. <https://doi.org/10.1177/1461444817706684>
- Cools, H., & Koliska, M. (2024). News automation and algorithmic transparency in the newsroom: The case of The Washington Post. *Journalism Studies*, 25(6), 662–680. <https://doi.org/10.1080/1461670X.2024.2326636>
- de Haan, Y., van den Berg, E., Goutier, N., Kruijckemeier, S., & Lecheler, S. (2022). Invisible friend or foe? How journalists use and perceive algorithmic-driven tools in their research process. *Digital Journalism*, 10(10), 1775–1793. <https://doi.org/10.1080/21670811.2022.2027798>
- Helberger, N., van Drunen, M., Moeller, J., Vrijenhoek, S., & Eskens, S. (2022). Towards a normative perspective on journalistic AI: Embracing the messy reality of normative ideals. *Digital Journalism*, 10(10), 1605–1626. <https://doi.org/10.1080/21670811.2022.2152195>
- Jia, H., Appelman, A., Wu, M., & Bien-Aimé, S. (2024). News bylines and perceived AI authorship: Effects on source and message credibility. *Computers in Human Behavior: Artificial Humans*, 2(2), Article 100093. <https://doi.org/10.1016/j.chbah.2024.100093>
- Jones, B., Jones, R., & Luger, E. (2022). AI “everywhere and nowhere”: Addressing the AI intelligibility problem in public service journalism. *Digital Journalism*, 10(10), 1731–1755. <https://doi.org/10.1080/21670811.2022.2145328>
- Lewis, S. C., Guzman, A. L., & Schmidt, T. R. (2019). Automation, journalism, and human–machine communication: Rethinking roles and relationships of humans and machines in news. *Digital Journalism*, 7(4), 409–427. <https://doi.org/10.1080/21670811.2019.1577147>
- Lewis, S. C., Sanders, A. K., & Carmody, C. (2019). Libel by algorithm? Automated journalism and the threat of legal liability. *Journalism & Mass Communication Quarterly*, 96(1), 60–81. <https://doi.org/10.1177/1077699018755983>
- Liu, B., & Wei, L. (2019). Machine authorship in situ: Effect of news organization and news genre on news credibility. *Digital Journalism*, 7(5), 635–657. <https://doi.org/10.1080/21670811.2018.1510740>
- Milosavljević, M., & Vobič, I. (2019). Human still in the loop: Editors reconsider the ideals of professional journalism through automation. *Digital Journalism*, 7(8), 1098–1116. <https://doi.org/10.1080/21670811.2019.1601576>
- Munoriyarwa, A., Chiumbu, S., & Motsaathebe, G. (2023). Artificial intelligence practices in everyday news production: The case of South Africa’s mainstream newsrooms. *Journalism Practice*, 17(7), 1374–1392. <https://doi.org/10.1080/17512786.2021.1984976>
- Ndlovu, M. (2024). Audience perceptions of AI-driven news presenters: A case of “Alice” in Zimbabwe. *Media, Culture & Society*, 46(8), 1692–1706. <https://doi.org/10.1177/01634437241270982>
- Porlezza, C. (2023). Promoting responsible AI: A European perspective on the governance of artificial intelligence in media and journalism. *Communications: The European Journal of Communication Research*, 48(3), 370–394. <https://doi.org/10.1515/commun-2022-0091>
- Raemy, P., Bendahan Bitton, D., Ötting, H., Hoffmann, C. P., Godulla, A., & Puppis, M. (2025). Deepfakes and journalism: Normative considerations and implications. *Journalism Studies*, 26(14), 1724–1744. <https://doi.org/10.1080/1461670X.2025.2547300>
- Thurman, N., Stares, S., & Koliska, M. (2025). Audience evaluations of news videos made with various levels of automation: A population-based survey experiment. *Journalism*, 26(1), 3–23. <https://doi.org/10.1177/14648849241243189>
- Toff, B., & Simon, F. M. (2025). “Or they could just not use it?”: The dilemma of AI disclosure for audience trust in news. *The International Journal of Press/Politics*, 30(4), 881–903. <https://doi.org/10.1177/19401612241308697>
- Valenzuela, S., Bachmann, I., Borah, P., & Solís Valdés, N. (2026). The effects of generative AI in news on media credibility and selectivity: Evidence from a conjoint experiment in Chile. *Digital Journalism*. Advance

online publication. <https://doi.org/10.1080/21670811.2026.2649018>

- Wang, S., & Huang, G. (2024). The impact of machine authorship on news audience perceptions: A meta-analysis of experimental studies. *Communication Research*, 51(7), 815–842. <https://doi.org/10.1177/00936502241229794>
- Wölker, A., & Powell, T. E. (2021). Algorithms in the newsroom? News readers' perceived credibility and selection of automated journalism. *Journalism*, 22(1), 86–103. <https://doi.org/10.1177/1464884918757072>
- Wu, S., Tandoc, E. C., Jr., & Salmon, C. T. (2019). Journalism reconfigured: Assessing human–machine relations and the autonomous power of automation in news production. *Journalism Studies*, 20(10), 1440–1457. <https://doi.org/10.1080/1461670X.2018.1521299>
- Zheng, Y., Zhong, B., & Yang, F. (2018). When algorithms meet journalism: The user perception to automated news in a cross-cultural context. *Computers in Human Behavior*, 86, 266–275. <https://doi.org/10.1016/j.chb.2018.04.046>