



"Transparency is More Than a Label": Audiences' Information Needs for AI Use Disclosures in News

Hannes Cools, Sophie Morosoli , Laurens Naudts , Karthikeya Venkatraj ,
Claes de Vreese & Natali Helberger

To cite this article: Hannes Cools, Sophie Morosoli , Laurens Naudts , Karthikeya Venkatraj ,
Claes de Vreese & Natali Helberger (01 Aug 2026): "Transparency is More Than a Label":
Audiences' Information Needs for AI Use Disclosures in News, Digital Journalism, DOI:
[10.1080/21670811.2026.2700592](https://doi.org/10.1080/21670811.2026.2700592)

To link to this article: <https://doi.org/10.1080/21670811.2026.2700592>



© 2026 The Author(s). Published by Informa
UK Limited, trading as Taylor & Francis
Group



View supplementary material [↗](#)



Published online: 01 Aug 2026.



Submit your article to this journal [↗](#)



Article views: 316



View related articles [↗](#)



View Crossmark data [↗](#)

“Transparency is More Than a Label”: Audiences’ Information Needs for AI Use Disclosures in News

Hannes Cools^a , Sophie Morosoli^b, Laurens Naudts^c,
Karthikeya Venkatraj^d, Claes de Vreese^e and Natali Helberger^f

^aAssistant Professor, ASCoR, Amsterdam, The Netherlands; ^bPostdoctoral Researcher, ASCoR, Amsterdam, The Netherlands; ^cPostdoctoral Researcher, IViR, Amsterdam, The Netherlands, Research Fellow, KU Leuven CITIP, Leuven, Belgium; ^dResearch Assistant, CWI, Amsterdam, The Netherlands; ^eUniversity Professor, ASCoR, Amsterdam, The Netherlands; ^fUniversity Professor, IViR, Amsterdam, The Netherlands

ABSTRACT

This study evaluates the role of disclosure transparency in rebuilding trust in journalism amid the increasing integration of generative AI (GenAI) tools in news production. While AI technologies enhance journalistic workflows by automating and augmenting content creation, they also introduce opacity in professional decisions, complicating traditional disclosure transparency norms and its practices. The paper maps citizens’ perceptions regarding disclosure transparency of AI use in journalism, addressing two key themes: (1) the practical information needs of news consumers about AI-generated content and (2) the specific rationales when exposed to (disclosures of) AI-generated news content. Through qualitative focus groups ($N=21$), the exploratory research reveals a strong demand for clear, visible, and detailed disclosures about AI-generated content. Participants emphasized the necessity of including general source references akin to traditional authorship attributions, explicitly stating AI involvement - for example, a label such as ‘generated by AI’ alongside author and publication details. Visual indicators like logos or watermarks in contrasting colors were preferred to ensure AI disclosures are noticeable and not easily overlooked. This study contributes to developing more granular audience-centered, practical guidelines for AI transparency in journalism that goes beyond the mere ‘label’, emphasizing that effective disclosure transparency requires more than simply ‘informing’ audiences.

KEYWORDS

Artificial intelligence;
audience perceptions;
disclosures; labeling;
responsible AI;
transparency

Introduction

The rapid advancement of Generative AI tools (GenAI) such as ChatGPT, Copilot, and DALL-E has already altered journalistic workflows, influencing every stage of news reporting – from gathering and production to distribution (Diakopoulos et al. 2024).

CONTACT Hannes Cools  h.cools@uva.nl  Assistant Professor, ASCoR, Amsterdam, The Netherlands

 Supplemental data for this article can be accessed online at <https://doi.org/10.1080/21670811.2026.2700592>.

© 2026 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group

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

These technologies enable newsrooms to automate tasks, enhance content creation, and streamline processes, but their integration has also sparked debate about ethical implications. As AI tools become more pervasive in journalism, the need for clear and consistent disclosure about their use has grown increasingly urgent (Stenbom, Wiggberg, and Norlund 2023). Central to this debate is 'disclosure transparency,' understood here as the obligation to disclose how AI tools are employed in the news production process so that audiences are informed and journalistic ethics are upheld (Diakopoulos 2015). Following Karlsson (2020), we use disclosure transparency as the sensitizing concept and do not address participatory or ambient transparency, which emphasize audience involvement (Karlsson 2010). Throughout the paper, the terms 'transparency,' 'disclosure transparency,' and 'disclosures' are used interchangeably; 'information needs' refers to the practical information audiences require to understand and evaluate the use of AI (Morosoli et al. 2025).

The importance of disclosure transparency in AI use has also become a focal point in regulatory frameworks regarding AI worldwide. For instance, the 'EU AI Act,' the 'US AI Bill of Rights,' and 'UNESCO's Ethics on AI' all emphasize the need for transparency in AI applications in relation to audiences (EU AI Act, 2025; UNESCO, 2022; US AI Bill of Rights, 2024). In this context, transparency is seen as a means to empower users, allowing them to make informed decisions and protect their interests (Zier and Diakopoulos 2024). Regulators and policymakers therefore prescribe transparency and the kinds of information the audience should receive. More often than not, however, these information requirements are not informed by empirical research into what kind of information audiences actually require in order to make informed choices, and how disclosure transparency can trigger desirable follow-up action, such as making more informed choices (Morosoli et al. 2025).

While these regulations might not target news organizations in a direct way, their principles have influenced the development of guidelines within the news industry (Cools and Diakopoulos 2023). For instance, many newsrooms have embraced disclosure transparency as a core value in 'responsibilizing' AI use, incorporating it into their ethical frameworks and operational practices. However, the practical implementation of these guidelines – such as the form, content, and frequency of AI disclosures – remains an underexplored area of research, particularly when examining how disclosure requirements are related to audiences' practical information needs (Toff and Simon 2024).

Although earlier research has investigated the specific uses of (Gen)AI (Cools and Diakopoulos 2026; d'Haeseleer et al. 2025), and the guideline documents related to its adoption (Becker, Simon, and Crum 2024; de Lima-Santos et al. 2025), it was only recently that scholars started to focus on how these disclosure requirements could be effectuated in practice. For instance, Altay and Gilardi (2024) have shown that labeling headlines as 'AI-generated' can reduce audiences' perceived accuracy and sharing intention. Relatedly, some recent studies have focused on the trade-off between transparency and trust perceptions of audiences related to AI disclosures in news (e.g., Gamage et al. 2025; Longoni et al. 2022; Prajod et al. 2026). Prajod et al. (2026), for example, found that audiences preferred detailed disclosures because they provided more transparency but also that these disclosures result in lower trust scores. Much will therefore depend on the information provided, and how well it succeeds

in signaling not only that a piece of content has been AI-generated, but also what has been done to make sure the content is trustworthy (Gamage et al. 2025). Similarly, Mattis et al. (2025) found that when journalists disclose the use of AI in news, trust in it is lower, a pattern also characterized in prior literature as the ‘transparency dilemma’ (Schilke and Reimann 2025). This dilemma underscores a deeper ambiguity in how audiences interpret and perceive trust in AI, as well as how they evaluate AI disclosures in news contexts. Despite these early insights, the existing research remains limited in scope and depth and often evaluates disclosure of AI use through binary frameworks - either disclosed or not disclosed and relies on quantitative methods. In particular, we know little about why specific audiences react positively or negatively to (labeling) AI-generated news (Karlsson 2020). Moreover, most studies focus on the effects of disclosure transparency on the audience but often overlook what specific information they need to be able to assess the use of AI in news.

In doing so, this study aims to add to this emerging field by exploring two central themes:

1. Audiences’ information needs in disclosing the use of AI in news.
2. Audiences’ perceptions of trust when exposed to (disclosures of) AI-generated news content.

By focusing on these themes, this study actively contributes to scholarly work on how audiences understand specific disclosures of AI use in journalism. As it seeks to explore information needs, and specific rationales, this study adopts a qualitative approach by conducting focus groups ($N=21$). With the increased adoption of (Gen) AI in newsrooms, it has become crucial to include and to understand audiences’ information needs, as it can help news organizations communicate their use of AI technologies more effectively and meaningfully while maintaining the existing trust relationship and helping the audience to make informed choices. This study shows that audiences are not only concerned with whether AI is used but also how it is used and disclosed. Ultimately, we believe that this study contributes to understanding specific disclosure information needs of individuals in a more granular way which could, in turn, result in a potentially improved trust relationship between news organizations on the one hand and news consumers on the other.

Literature Review

Newsroom’s Use of (Generative) AI

The integration of AI in journalism is reshaping the industry by (partially) altering professional routines and skills, ultimately changing the nature of journalistic outputs (Simon 2024). As these emerging technologies continue to develop and evolve, they present potential tensions with the journalism ethics framework, particularly in how they might influence these routines and skills, and to what extent journalists and news organizations rely on such systems. For example, AI may unintentionally introduce bias at various stages across the journalistic value chain that can lead to discrimination (Caswell 2023), both in term of amplifying already-existent stereotypes

and gendered language, as well as what is being deemed newsworthy. As AI has already demonstrated the ability to generate text, images, and videos from short prompts or even single words, it raises additional concerns about the accuracy, authenticity, and credibility of journalistic content (Beckett and Yaseen 2023; Jones, Luger, and Jones 2023).

Despite the challenges that (generative) AI tools bring about, AI also comes with many opportunities for the media, and their relationship with audiences. Established news organizations are investing in and implementing various technologies like Large Language Models (LLMs) to help generate and augment news articles, such as financial summaries and sports recaps. For example, *The New York Times* introduced AI tools in 2025 for tasks such as editing, summarizing, coding, and writing, aiming to enhance efficiency while maintaining journalistic integrity. *BBC News* is establishing a new department focused on integrating AI to offer more personalized content, particularly targeting younger audiences who consume news *via* smartphones and social media platforms.

As (Gen)AI and its tools are increasingly being implemented in newsrooms, there are new ethical considerations that have sometimes led to guidelines for the ethical use of AI in newsrooms that go beyond the current journalistic ‘charters’ that many newsrooms have (Cools and Diakopoulos 2023). Prominent ethical principles regarding the use of AI include the ‘human in the loop mechanism’, ‘upholding journalistic quality standards’ and ‘disclosure transparency requirement’ (e.g., Becker, Simon, and Crum 2024). The last principle, namely the ‘disclosure transparency requirement’ is especially dominant when deploying AI in journalism for at least two reasons. First, disclosure transparency could (Karlsson 2010) maintain trust between news organizations and their audiences by ensuring that readers understand how AI-generated content is produced and verified (Toff and Simon 2024). Second, transparency has been a core ethical concept for the news industry for decades and a valuable norm to study the deployment of AI in newsrooms, especially because different forms of transparency could foreground audiences’ trust perceptions (Karlsson 2010; Koliska 2022). In the next section, transparency is discussed as this core feature of journalistic practice, as well as a general, and often vague principle that is front-and-center in current regulatory frameworks on AI (Cools and Diakopoulos 2026).

Transparency: A Journalistic Norm and Practice in the Age of AI

Over the past decades, research has conceptualized transparency not merely as a means of accountability but as a symbolic and ritualistic practice that signals journalistic credibility and a professional identity to the public (Deuze, 2005; Karlsson 2010). Scholars have also traced its deeper foundations – examining, for example, how philosophical debates about truth and openness, societal expectations of institutional accountability, the historical evolution of press norms, and the influence of digital technologies have collectively shaped transparency as both an ideal and a practical concern in journalism (Plaisance, 2007; Vos and Craft, 2017). In doing so, both conceptual and empirical contributions in *Digital Journalism* and beyond have resulted in a more granular understanding of what the norm entails (e.g., Diakopoulos 2015; Heim, 2024; Zamith, 2019). Because of these contributions, scholarly work no

Table 1. Forms of transparency (building on Plaisance, 2007; Karlsson 2020; 2021).

	Disclosure transparency	Participatory transparency	Ambient transparency
Conceptualization	Forms of transparency where news organizations disclose, explain and are open about the way news is gathered, produced and distributed.	Forms of transparency where news organizations invite users to (actively) participate in different stages of the news production process.	Forms of transparency where news organizations make it possible for news consumers to evaluate and form meaning of news without explicit explanation.
General examples	Explaining errors; justifications; general disclosure on the production process.	Crowdsourcing information; embedded audience content from social media; eyewitness reports.	Hyperlinks; profiles of journalists.
Examples in the context of AI	LLM provider; code; prompt information; specific usage of AI in article.	Statements on usage of user data in relation to AI, forms	Links to professional guidelines on AI; terms of use (ToU); access to data policies.

longer defines transparency as merely a way of holding journalism accountable but evaluates the concept as an active practice that operates at several levels. To bring structure to this literature, we build on the scholarly work of Karlsson (2020; 2021) and Plaisance (2007) and highlight three more ‘practice-oriented’ and conceptually interconnected forms of the concept (see Table 1). We also highlight how these three forms of transparency could be practiced in the age of AI, as literature on these practiced forms of transparency has been fragmented (Cools, 2025).

First, disclosure transparency in journalism refers to explicit, story-bound revelations about how and why a news item was produced, such as explaining why a story was published, how it was sourced and framed, and when and why it was updated (Karlsson 2010; 2020). This form of transparency has been operationalized and practiced through a set of ritualized transparency techniques like detailing the production process, admitting errors, linking to original material, showing that audiences tend to see such explanations as an important aspect of ‘open’ journalism. In the context of the disclosure of AI use in journalism, information that could be disclosed here includes, but is not limited to the specific LLM provider, code and prompt information (Bengani, 2022; citation withheld). Second, *participatory transparency* focuses on giving opportunities to audiences to contribute or provide (general) feedback to newswork (Karlsson 2020). This form emerged alongside the rise of digital platforms that enabled two-way communication between journalists and audiences, transforming traditional gatekeeping practices (Singer et al. 2011). By inviting reader comments, crowdsourcing information, or hosting public forums, news organizations signal openness while simultaneously negotiating the boundaries of professional authority (Westlund and Lewis 2014). In the context of the disclosure of AI use in journalism, participatory transparency could include information to audiences about the role of their (personal) data. Third, *ambient transparency* that centers around contextual cues around stories (e.g., hyperlinks, journalists’ profiles/opinions) that help audiences evaluate content without explaining it directly (Karlsson 2020; Zamith, 2024). Ambient cues are often embedded in site architecture or platform interfaces rather than individual articles, and can signal, for example, editorial lines, source diversity, or connections to related news coverage (Fletcher and Nielsen 2024) In the context of the disclosure of AI use

in journalism, ambient forms of transparency could include professional and organizational guidelines on AI, Terms of Use (ToU), and access to related data policies.

In this study, and building on the work of Plaisance (2007) and Karlsson (2020, 2021) (Table 1), we focus on disclosure transparency as our core sensitizing concept, and we define it as the obligation to disclose information about how AI tools are employed in the news production process, ensuring that audiences are informed and that ethical journalistic principles are upheld. From 2022 onwards, generative AI technologies have entered news organizations and have led to active experimentation in their production processes to gather, produce and distribute news. However, there is considerable variation in whether and how these news outlets practice disclosure transparency of AI use to their audiences, making these disclosures often inconsistent, vague, or hidden deep within articles or on their websites (Morosoli et al. 2025; Toff and Simon 2024). This complicates audiences' understanding of AI use in news, and this understanding matters, as research has shown that (mis)interpretation of such AI disclosures directly impacts public trust: if readers cannot distinguish between AI-generated and human-created content, they may feel deceived, which undermines confidence in the news (Gilardi and Altay, 2024). Similarly, the effectiveness of transparency and disclosure practices remains somewhat an open question, especially in relation to their audiences, as studies suggest that audiences may not always engage with or fully understand AI disclosures when they are provided (Gilardi and Altay, 2024).

Within AI legislation, transparency is frequently framed as a foundational tool for audiences' empowerment and informed choice. However, for transparency to fulfill this function meaningfully, informational obligations must directly address the specific risks audiences face when engaging with AI-generated media content – such as exposure to misinformation, manipulation, or discriminatory outputs (Piasecki et al. 2024), and the specific information needs audiences have to be able to navigate these risks. This also implies that disclosure policies should not only inform but also be actionable, in the sense that they equip audiences with the knowledge and tools needed to act on that information. In this sense, transparency becomes a necessary precondition for citizens' empowerment, rather than a goal in itself (Morosoli et al. 2025).

However, in practice, tools that enhance user agency, such as actionable rights, meaningful alternatives to AI-generated news, or mechanisms to contest editorial decisions involving AI, are often absent or underdeveloped (Piasecki et al. 2024). In such contexts, transparency risks serving only a dignitarian function: namely superficially signaling respect for the public without enabling real intervention or accountability (Piasecki et al. 2024). At the same time, legal transparency obligations can come into tension with media independence. Under the EU's AI Act, for instance, the disclosure of AI-generated or manipulated textual, public interest content is therefore not required if the material has been reviewed or published under human editorial control (Article 50, para. 4 AI Act), effectively leaving disclosure decisions to news organizations themselves. As governance frameworks for AI in news organizations continue to evolve, there is growing value in understanding what kinds of information audiences actually need, not just to be informed, but to pursue meaningful follow-up actions in response to AI's role in news production.

The Audience Perspective on Disclosing AI Use in Journalism

The relationship between transparency and trust is central to understanding how audiences respond to AI disclosures in journalism, because trust underpins credibility, perceived legitimacy, and downstream behaviors such as selection, sharing, and subscription (Diakopoulos 2015). Disclosure transparency could be a prerequisite for news organizations to calibrate the adoption of AI; miscalibrated disclosure can backfire, eroding trust and perceived integrity, so identifying when disclosure transparency helps or harms is essential for newsroom principles and practices.

While some research suggests that disclosure transparency can enhance credibility (Jung et al. 2026), other studies report negligible effects or even diminished trust (Altay and Gilardi, 2024; Graefe et al. 2018). Existing scholarship suggests that disclosure transparency's impact on trust could depend on specific conditions. First, the length and detail of disclosure could matter: Transparent practices that demonstrate editorial oversight and ethical reasoning tend to build trust, whereas vague or technical disclosures may confuse audiences and undermine credibility (Wölker and Powell 2021; Diakopoulos et al. 2024). More concretely, 'proportional disclosures', namely more lengthy statements that specify what AI did and did not do, could elucidate the amount of human oversight, which could foster trust among engaged audiences (Diakopoulos 2015; Prajod et al. 2026; Wölker and Powell 2021). By contrast, vague, purely technical, or boilerplate disclosures or simple labels often fail to provide diagnostic information and can reduce perceived trustworthiness (Prajod et al. 2026).

Second, contextualization and tone might shape trust outcomes: when AI is presented as enhancing journalistic quality, audiences respond more favorably than when it appears to represent cost-cutting or reduced human judgment (Gamage et al. 2025; Beckett 2019). Experimental work on AI attribution suggests that signaling AI involvement can shift perceived bias and credibility, depending on how the attribution is presented and contextualized (Waddell, 2026). Taken together with labeling research (Altay and Gilardi, 2024), these findings indicate that bare 'AI-generated' tags can trigger negative heuristics, whereas contextualization might mitigate such effects.

Third, digital literacy of audiences could impact trust of disclosing the use of AI in news: engaged, critical news consumers could scrutinize AI disclosures more intensely and demand further explanation, while habitual news users may remain largely indifferent to production details (Swart et al. 2022). For example, engaged users with higher levels of digital literacy users might seek specificity about AI use in news, resulting in greater trust while penalizing vague or symbolic transparency (Ahva, Heikkilä, and Kunelius 2015; Diakopoulos 2015; Wölker and Powell 2021).

Overall, scholarly work has shown that audiences have expressed a strong desire for disclosure transparency in news. Newman et al. (2024) argue that the desire is rooted in concerns about manipulation, the erosion of editorial integrity, and the audience's ability to manage their trust relationship with media outlets. Studies show that many individuals feel manipulated if they are not informed about the use of AI in news production (Piasecki et al. 2024). Altay and Gilardi (2024) find that users in the U.S. and U.K. strongly support mandatory AI labeling, particularly to guard against

misinformation. At the same time, the design and placement of disclosure transparency do matter for audiences' trust perceptions (El Ali et al. 2024; Venkatraj et al. 2025). A simple 'AI-generated' tag may confuse readers who lack media literacy or lead them to distrust such content outright, even if it is accurate (Zier and Diakopoulos 2024). In this study, we therefore explore the specific needs audiences associate with AI disclosures in journalism, focusing on how disclosure transparency can serve not just as a symbolic gesture, but as a meaningful tool for informed news consumption. This leads to the following research question:

RQ1: What are audiences' information needs on how the use of AI in news is disclosed?

As mentioned above, the relationship between the explicit disclosure of AI use and audience trust is complex and context dependent. Disclosure transparency can indeed enhance trustworthiness and credibility perceptions among audiences (Jung et al. 2026), other studies report negligible effects or even diminished trust, as disclosures may heighten skepticism (Altay & Gilardi, 2024). Such conflicting evidence highlights the importance of examining how different types and timings of disclosures shape audience engagement. For instance, post-hoc disclosures (revealed after content consumption) have been shown to trigger 'retroactive distrust,' whereas pre-emptive disclosures may encourage audiences to engage more critically with the news (Ananny and Crawford, 2018). Taken together, this dissensus in scholarly work points to a gap in understanding how AI use and its disclosure practices influence trust, motivating the second research question:

RQ2: How do (disclosures of) AI-generated news content influence audience perceptions of trust to engage with the news?

Method

Focus groups were selected as the most appropriate method for this study because our research questions require understanding not only what information audiences need about AI-generated content, but also how they collectively make sense of AI disclosure practices and navigate trust decisions in response to such content. Unlike individual interviews or surveys, focus groups enable observation of interactive sense-making processes, how participants articulate, negotiate, and refine their concerns through dialogue with others (Kitzinger 1995). This method also allows researchers to probe emerging themes in real time and observe how group dynamics influence the development of shared understandings about appropriate transparency practices (Morgan 1997). These interactive processes are particularly valuable for exploring normative questions about what information should be disclosed about AI use in journalism, as participants can challenge each other's assumptions and collectively articulate standards that might not emerge through individual reflection alone.

To assess individuals' information needs and follow-up behavior when exposed to AI-generated news, we conducted three focus groups with Dutch citizens ($N=21$). The participants were recruited with the help of a local panel company (Bureau Fris)

and the focus groups were conducted in person at the authors' university. The participants received an incentive directly from the panel company after having taken part in the focus groups. Although focus groups are not intended to produce statistically representative samples, we aimed to approximate a diverse cross-section of participants to ensure a broad range of perspectives relevant to the study. Each group consisted of seven individuals, diverse in age, gender, and education levels. The youngest participant was 24 years old, and the oldest was 63 years old. Eleven participants identified as female and ten as male. Eight individuals had a lower education level, eight a medium, and five a high level of education. Four individuals were students, three were without an occupation, and fourteen were part of the workforce. The focus groups lasted approximately one hour and were conducted by the authors in the summer of 2024. In each focus group, two authors were present, one communication scientist and one legal scholar, with one leading the focus group while the other took notes. The legal scholar was included to provide expertise on AI-related legal obligations, which were also discussed in the second part of the focus groups. We chose in-person sessions to foster richer interactional cues (e.g., turn-taking, alignment, disagreement) that aid interpretation of disclosure transparency perceptions.

Focus Group Structure

The focus groups were structured in two parts. First, all focus groups began with the interviewer asking participants to write down their three main concerns about AI-generated content. This individual reflection phase ensured that all participants had the opportunity to formulate initial thoughts before group discussion, preventing dominant voices from prematurely shaping others' views. These concerns were then ranked by the group from least to most concerning. This ranking exercise required participants to negotiate priorities, articulate reasoning for their choices, and respond to others' perspectives, generating the kind of interactive dialogue central to focus group methodology (Morosoli et al. 2025). The deliberative process revealed not just what participants were concerned about, but why certain concerns resonated across the group while others were debated or deprioritized. Every individual received the same examples. The AI-generated images and text were sourced online from websites that were flagged as 'AI-generated news' by Newsguard¹ on topics of politics, sports and art. Participants were informed that all the material was generated by AI and given as much time as needed to review it. The headers included various news topics, containing either only AI-generated text or a headline and an AI-generated image. Two examples were used in the focus groups (See Appendix A). First, the political header example about Volodymyr Zelenskyy buying a yacht with money given to Ukraine from the US government. Second, the sport related AI-generated news article about how to play like a professional volleyball player. After reviewing the material, participants engaged in a plenary discussion about their information needs, expectation management, need for rights, and the authenticity of AI-generated content. This study focuses specifically on responses related to individuals' information needs and individuals' follow-up behavior after seeing AI-generated news content. In addition, the inclusion of concrete AI-generated news examples served as shared reference

points that grounded abstract discussions about transparency in participants' lived responses to actual content. This approach differs from quasi-experimental designs: rather than measuring individual-level effects of specific disclosure formats, we used the examples to give participants the idea of what AI-generated news could look like (For a more detailed version of the focus group protocol, see Appendix B).

The focus groups were recorded, transcribed, and translated by four of the authors. To protect participant identities, we anonymized the data using 'Subject 1, Group 1' as a reference. For readability and clarity of quoted materials, filler words and stammering were edited out.

For the analysis, we adopted an inductive approach using thematic analysis to identify recurring patterns (Braun & Clarke, 2006). This method is known for its flexibility in sample size, data collection, and theoretical flexibility. The authors first read the transcripts multiple times to become familiar with the content, and each created an individual codebook using open coding to identify themes and possible codes. Based on these codebooks, a consensus codebook was created with all authors present (see Appendix C). The authors then analyzed all transcripts line-by-line together, discussing and coding each statement using NVivo software. Throughout the coding process, codes and themes were thoroughly and regularly discussed by all authors to ensure validity. This "negotiated agreement" process was relied on to ensure inter-coder reliability (Campbell et al. 2013).

Country Selection

The Netherlands presents a particularly interesting case for studying artificial intelligence in journalism due to its high level of digitalization, strong public media tradition, and relatively high trust in news compared to other Western countries (Newman et al. 2024). Dutch news organizations have been early adopters of digital innovation, including algorithmic tools and AI applications in editorial workflows (Cools and Diakopoulos 2026). Moreover, the Dutch media landscape is characterized by commercial outlets like *Mediahuis* and *DPG Media*, which are conglomerates of both national and regional titles in the Dutch Belgian media landscape (Commissariaat voor de Media 2024). There is also a relatively strong public broadcaster *NPO* which is generally trusted by around 80% of the Dutch population, according to the Digital News Report (Newman et al. 2024).

Results

The results are described in two main parts: (1) audiences' information needs on how AI use is disclosed and (2) how disclosures of AI-generated news content influence audience perceptions of trust. Existing scholarly work is included in the findings section to contextualize the results within the broader field of *Digital Journalism*.

Audiences' Information Needs on How AI Use is Disclosed (RQ1)

First and foremost, the focus groups revealed a shared expectation that AI-generated journalistic content should be disclosed. One suggestion was the inclusion of a general

statement, similar to traditional attributions of authorship, but explicitly stating something along the lines of 'AI was used'. A participant noted that articles typically include details like the author's name and publication date, and AI-generated content should follow the same practice:

"A source reference does seem useful. Normally, you have something like an author, date and place. And then it should say, something like, 'generated by AI'" (Subject 5, Group 1)

This argument is also underscored in the report by Nordic AI journalism, which stipulates that disclosure will allow audiences to make informed decisions about the validity and credibility of the information, and that it should be mentioned clearly on the level of the article (Nordic AI, 2024). This aligns with prior research showing audiences value transparency as a mechanism for accountability in digital journalism (Cools 2025; Vos and Craft, 2017). However, unlike earlier studies that mainly addressed human authorship and sourcing, our findings extend this discussion by focusing on algorithmic authorship and specifying the granular disclosure forms audiences expect.

Another theme that was discussed in the focus groups had to do with the importance of visual indicators to make AI disclosures more noticeable. While some participants preferred a small but clear label, others stressed that it must be highly visible and distinct from other elements on the page. As a participant argued, *"a logo or watermark should stand out in contrasting colors to avoid being overlooked"* (Subject 5, Group 2). Others, like Subject 1 and Subject 5 from Group 2, suggested either a 'standardized certification' or a 'seal of approval'. These suggestions reflect what other scholarly work has identified as the need for standardized visual cues that can become part of audiences' cognitive processing of news content (Strikovic and Cools 2025). Similarly, Trattner et al. (2026) demonstrate that standardized provenance labels such as those developed by the C2PA (Coalition for Content Provenance and Authenticity) initiative increase trust in news platforms across Western countries, reinforcing our participants' call for highly visible and consistent disclosure practices.

As concerns were raised about disclosures being hidden in footnotes or at the end of an article, participants also reflected on the placement of AI disclosures. Participants were advocating for positioning transparency labels at the very beginning of an article to ensure immediate transparency:

"Maybe at the beginning of the text it says sometimes 'Amsterdam' right? That it then says, 'AI-generated'" (Subject 4, Group 1)

Other participants believed a more integrated approach would be effective, such as placing a small but recognizable label throughout the article to maintain awareness as the reader progresses. Overall, there participants mentioned that AI disclosures should not be hidden or difficult to spot. At the same time, our results resonate with Toff and Simon (2024), who caution that disclosures can create dilemmas: while audiences demand them, they can also trigger reduced trust if the disclosure is perceived as vague or poorly contextualized.

Participants discussed the varying levels of disclosure of AI use and whether a simple 'AI-generated' label adequately reflects reality. Some noted that AI is often

used partially – for tasks like translation, summarization, or editing – rather than producing entire articles from scratch. A participant suggested a proportional disclosure system with different levels, distinguishing between ‘fully AI-generated content’, ‘partially AI-assisted work’, and ‘low AI-augmented material’ (Subject 2, Group 3). This proportional form of transparency would allow readers to understand the exact role AI played in the content’s creation, and would provide greater granularity to the disclosure of AI use.

Connected to this, participants also made a case for use, topic, and news outlet differentiation. For use differentiation, participants expressed varying views on the acceptable uses of AI in journalism, with a clear distinction between AI-assisted editing and fully AI-generated news. Some found AI’s role in photo generation particularly sensitive, as ‘fabricated’ images could mislead audiences or misrepresent individuals. A participant emphasized this concern, stating, *“Imagine: They make an article about you, you’re suddenly famous, and then suddenly you’re in a photo, and it’s not you. That, I think, that’s very helpful for the person being written about, that it then says ‘AI’, that everyone knows it’s not real.”* (Subject 4, Group 1). Others distinguished between AI being used as a supportive tool, such as refining or fact-checking text, and AI autonomously creating news stories. For instance, Subject 1 noted that while using AI for tasks like polishing text or verifying sources might not require disclosure, entirely AI-generated articles demand explicit disclosure:

“Actually, generating news, saying: make an article about Zelensky, what’s the latest news? And just posting that right away, that’s something else.” (Subject 1, Group 1).

For topic differentiation, participants drew contrasts between politically sensitive reporting and lighter subject areas. They described AI-generated news about global political figures, such as Zelensky or Biden, as requiring more extensive and explicit transparency than sports or entertainment coverage. This aligns with prior research showing that political reporting is particularly sensitive to transparency cues, and more recent work (Toff and Simon 2024) highlights how disclosure in high-stakes topics can backfire if it raises doubts about editorial rigor rather than informing or reassuring audiences. One participant questioned the necessity of disclosure in less critical reporting, stating: *“Say, one of those volleyball articles, yeah, that’s a bit informative about how popular the sport is. Then I would think, well, then I might believe it is AI. (...) But with something like that with Zelensky which is so timely and still so actively in the news every day, I would assume that less.”* (Subject 7, Group 1). This might indicate a ‘hierarchy of importance’, where AI use in politically sensitive topics was seen as more impactful and potentially problematic than its use in feature writing or lifestyle reporting.

The perceived credibility of AI-generated content was also tied to how participants viewed the news outlet that published it. So, for news outlet differentiation, participants expressed greater confidence in established media organizations adhering to ethical standards, making them more accepting of AI use in such outlets. Subject 7 from Group 2 mentioned that articles labeled as coming from *“ANP [largest Dutch news agency] or so... A lot of people already find that reliable”*, suggesting that AI-generated content might be more readily trusted if associated with well-established journalistic brands. However, there were also concerns about the generic nature of

AI-generated news in larger outlets, with one participant pointing out that articles from mainstream news organizations do not mention author information:

“If you get an article from NOS [public broadcaster in The Netherlands] or AD [commercial news outlet in The Netherlands], for example, you don’t necessarily know who wrote it.”
(Subject 5, Group 3)

This quote illustrates the tension between institutional trust and the absence of a named human author. Overall, the results of these practical requirements reveal a nuanced picture of how interviewees interpret AI use in journalism, in which disclosure expectations are closely linked to context, subject matter, and perceived trustworthiness of the news outlet (See [Table 2](#) for an overview). In addition, these findings also correspond to recent findings that institutional trust mediates content credibility (Schilke and Reimann 2025) but also resonates with Trattner et al. (2026), who show that provenance labeling interacts with source cues: established outlets benefit most from standardized transparency mechanisms.

Table 2. Practical information needs mentioned by participants.

Category	Description	Example
Logos and watermarks	Use of visible logos and watermarks to indicate AI-generated content.	An AI-generated image includes a small watermark in the corner reading ‘AI-generated’.
Contrasting colors and size	Use of contrasting colors to make disclosures of AI indicators stand out.	A yellow banner at the top of the page in bright contrasting color states ‘generated with AI support’.
Placement of indicators/ labels	Placement of AI indicators at the beginning of articles or prominently within the content.	A disclosure line directly under the headline reading: ‘This article contains AI-assisted reporting in the introduction.’
Use differentiation	Different uses require different forms of disclosure.	If AI is used only for grammar corrections, the disclosure reads ‘proofread with AI tools,’ whereas if text is generated, it says ‘This article has been generated by AI.’
Topic differentiation	Different topics require different forms of disclosure and different context.	For an election-related story, a detailed explanation of AI’s role is included, while for lifestyle pieces only a brief one-line disclosure is provided.
News Outlet differentiation	Different outlets require different forms of disclosure.	A more formal outlet uses standardized disclosure language (‘AI-assisted reporting’), while an entertainment-oriented outlet uses a visual icon to indicate AI use.

Despite an overall strong request for AI disclosures, participants also mentioned the ineffectiveness of disclosing the use of AI in news. One of the key issues discussed was the visibility of labels indicating AI-generated content. Subject 2 from Group 1 questioned whether small or inconspicuous labels would be noticed, asking, *“Well, if you see three here [refers to social media icons] and then once there’s a tiny little fourth [icon indicating the artificial nature of the content] or, then there’s a fourth logo, do you always see that, or do you have to go looking for it first? Like this is such a crazy article, where is there a logo?”*. This comment reflects the challenge of ensuring that labels are noticeable enough to inform readers without being overly intrusive.

Participants also questioned whether readers would actively attend to such indicators in everyday news use. Subject 5 from Group 3 noted, *“You almost wouldn’t pay*

attention to this, how small it is". Another concern was the ambiguity surrounding what AI-generated content actually means and how it affects the perceived reliability of the information. Subject 1 from Group 2 pointed out, *"Oh yes, that's possible, but I think ultimately, it's not necessarily about whether it's then generated by AI, you can put that in there, but I don't think it has that much to do per se with reliability at this point in time"*. This concern directly speaks to what scholars have called 'banner blindness', where users systematically ignore standardized interface elements that resemble advertisements or routine disclaimers (Dobber et al. 2025). Additionally, it might suggest that simply labeling content as AI-generated may not fully address readers' concerns about accuracy or trustworthiness. Subject 4 from Group 1 further elaborated on this issue, stating:

"If, at the top of an article, it would say 'written by AI', even though you copy-pasted it in an AI generator and asked: 'What should I and shouldn't I leave out? (...) Then it's actually also through AI, but that doesn't mean it immediately becomes less reliable, because it can also remain your own words'".

Additionally, some participants expressed concerns that labeling content as AI-generated could inadvertently invoke distrust. Subject 4 from Group 2 noted, *"So it can also scare people a little bit like: oh, am I even going to read it at all, if there is such an icon?"*. This suggests that labels indicating AI use might deter readers from engaging with the content altogether, regardless of its actual reliability. This potential for labels can therefore backfire and highlights the delicate balance between transparency and the risk of alienating audiences.

Finally, participants questioned the long-term relevance of labeling AI-generated content. Subject 2 from Group 3 speculated, *"If we look 20 years from now, when almost everything will soon be written with AI, then I think, that generation doesn't care so much anymore if there's an AI logo or something along those lines"*. This observation aligns with research on generational differences in AI acceptance, which suggests that digital natives may develop different expectations for AI transparency as artificial intelligence becomes increasingly integrated into information environments (d'Haeseleer et al. 2025). This temporal perspective raises important questions about the sustainability and effectiveness of current disclosure approaches, suggesting that transparency frameworks must be designed to evolve alongside changing technological and social contexts, also incorporating trust perceptions of audiences (Schilke and Reimann 2025).

Audiences' Perceptions of Trust When Exposed to (Disclosures of) AI-Generated News Content (RQ2)

Having identified existing information needs when it comes to disclosure transparency, it is important to investigate individuals' behavior when being confronted with AI-generated news content. Across the focus group discussions, participants repeatedly returned to feelings of skepticism. As Subject 2 from Group 1 noted, *"Well, you view the reporting somewhat skeptically to begin with"*. This sentiment was echoed by Subject 4 from Group 2, who stated, *"If I knew that it was all generated by AI, I would indeed look at it with a critical eye. Maybe I wouldn't take the news with a grain of salt but with*

a grain of truth?”. The phrase of ‘taking news with a grain of salt’ recurred across the different focus groups, indicating a tendency to approach AI-generated news with a degree of doubt. Subject 1 from Group 3 elaborated on this point, suggesting that the level of skepticism might vary depending on the stage of AI development:

“I think it depends on what stage of AI you're in. Because now everything that is generated is often not quite right yet, but so then you take it with a grain of salt. But I think that in, I don't know, 10 years, they'll probably be able to make something that does in fact make it very reliable”.

The skepticism towards AI-generated news is further compounded by the association of such content with ‘fake news’ and ‘clickbait’. Subject 4 from Group 1 highlighted this issue, stating, *“People who indeed offer that, so you also see it sometimes on Instagram, that such a piece comes between your story, and then it says for example ‘Eva Jinek fired’ [Dutch talkshow host], whatever, and then you think, what is that? But yes, that's also all very AI-generated and so that's fake news, just to get clicks. (...) So, that's why I have a negative image of AI when it's a news item and then I think: no, I'm not even going to read that then”.*

This reaction underscores that AI is also used to generate fake news which has the potential to erode trust in news sources, which was underscored in the focus groups. Additionally, it demonstrates what recent scholarly work has identified as ‘negative transfer effects,’ where experiences with low-quality AI applications influence perceptions of legitimate journalistic uses (Toff and Simon 2024). In this light, disclosing AI use can also lead to a loss of trust in previously reliable sources. Subject 3 from Group 1 stated:

“But the moment I would come across something similar from a source I trust, for example, my trust would be completely gone. Then I wouldn't be interested in other articles after that either”.

This quote suggests that the use of AI in news production, if not managed transparently and responsibly, may have a ‘backfire effect’ on the credibility of entire news outlets (Ananny and Crawford, 2018). Subject 4 from group 2 stated: *“I'm pro AI anyway! But reading a disclosure with those capital letters? It just doesn't read very well either. It feels robotic, so that's also why I would stop reading.”* Recent work by Trattner et al. (2026) shows that CP2A provenance labels can help restore trust at a time when the boundaries of what counts as ‘authentic’ and ‘truthful’ news are increasingly blurred. The disclosure debate has already intensified, with outlets such as Hoodline and CNET facing criticism for their lack of transparency regarding the extent of AI involvement in content production.

Discussion and Conclusion

This study explored audiences’ information needs and trust perceptions related to disclosure transparency about the use of AI in journalism. In addressing our first research question (RQ1) - what practical information audiences require regarding AI use - the focus group interviews indicate that participants place considerable value on transparency and disclosure. Specifically, they emphasized that AI labels should

be clear, recognizable, and trustworthy. Such transparency, particularly when backed by credible institutions, was seen as a means of strengthening trust between news organizations and their audiences. The specific content and form of the disclosure of AI also matter according to our participants. Specifically, they mention the placement, contrasting colors and size, as well as recognized logos and watermarks. These granular information needs highlight the importance of immediate and clear disclosure formats to prevent feelings of manipulation (i.e., fear of fake news) and to uphold the existing trust relationship between audiences and news outlets. Individuals also differentiate disclosures with respect to specific uses, within certain topics, and they also distinguish between news outlets, often in accordance with the trust relationship they already have with these outlets.

For the second research question (RQ2) – how AI-generated news content influences audience perceptions of trust to engage with the news – we found that AI disclosures often prompted skepticism and fueled distrust. From the perspective of transparency as a way to empower citizens and support informed decision-making, the results indicate that transparency can spark reflection and caution. At the same time, it may also influence citizens' trust and willingness to engage with the news. Moreover, participants often made trust-related judgments, both positive and negative, when the presence of AI was disclosed or inferred, before parsing the actual specifics of that disclosure. This finding could suggest that disclosure transparency frequently functions as a signifier that surfaces an underlying judgment about AI-mediated journalism.

This finding, together with the observation by our participants that the information that content is generated by or with AI does not yet provide sufficient information to assess the accuracy and trustworthiness of the content triggers the follow-up question of what information users would need to be able to decide whether or not AI-generated content can be trusted. It does, however, show that transparency is more than a label, and that more factors and features should be taken into account when considering information needs in relation to AI use in news. For example, they were more accepting of AI use, particularly for low-stakes tasks like proofreading. However, they demanded full disclosure for more substantive applications such as content creation or image generation. This pattern suggests that interviewees treat AI presence as a trust cue that depends on what the technology does in a given context: low-stakes presence can be accommodated with less concern, whereas high-stakes presence leads to closer scrutiny, with the quality of disclosure shaping how this scrutiny unfolds. This nuanced finding reflects the broader shift in journalism toward greater audience agency, where trust hinges not just on disclosure but also on how it is communicated (Fletcher and Nielsen 2024; Swart et al. 2022). At a broader level, this exploratory study underscores a core dilemma for news organizations: not only whether to disclose AI use, but also how to do so effectively. Based on our findings, three key implications and recommendations for scholarly work, and the news industry emerge:

1. Being transparent is not enough

While transparency about the use of AI in news production is essential to fostering trust, this study suggests that transparency alone is not enough. To be effective, disclosures must be designed with the audience in mind - clearly

visible, intuitively placed, and supported by recognizable indicators such as logos or labels (Zier and Diakopoulos 2024). This ensures that audiences can quickly grasp when and how AI is involved in content creation, aligning with both journalistic standards and emerging regulatory expectations like those outlined in the EU AI Act (Schnell, 2024).

2. The design and amount of information in the disclosure matter

Disclosures should go beyond simply stating that AI was used, putting both the design and the amount of information in focus. Participants in this study expressed a clear need for contextual information that helps them assess the reliability, accuracy, and editorial oversight of AI-generated content. In this way, transparency becomes not just a mechanism of disclosure, but a form of media education - helping audiences understand that the use of AI does not inherently diminish content quality or truthfulness, provided it is used responsibly and ethically (Morosoli et al. 2025).

3. Information needs of disclosures are individual and contextual

The findings of this study have shown that information needs of disclosures are not one-size-fits-all. Audiences differentiate between types of content (e.g., headlines, full articles, images) and the perceived risk or importance of the material influences how much and what kind of disclosure they expect. News organizations that already have high levels of public trust may have more flexibility in how they communicate about AI use, but clarity and openness remain vital.

Reflecting on these implications and recommendations, we argue that meaningful transparency involves engaging the public in how disclosure is implemented - inviting them into the design process and acknowledging that perceptions of trustworthiness are not universal but highly contextual. In this sense, AI transparency as a principle should become a co-constructed experience in practice, shaped by the audience's expectations, prior knowledge, and interaction with the content (see also: El Ali et al. 2024). This makes the challenge of effective AI disclosure less about technical compliance and more about relational trust, which results in considering 'the eye of the beholder'.

This study has limitations. Firstly, our research methodology was based on focus groups, which only offer a partial view of the broader audience landscape vis-à-vis AI disclosures in news headers. The timing of the interviews should also be considered when reading the results, as they only provide a snapshot of individual perceptions at the time of data collection. To garner a more comprehensive understanding across a wider sample of individuals in different contexts, we recommend supplementing this study with longitudinal research that evaluates the changing uses and implications of AI disclosures in journalism, also among different demographics. More concretely, experiments could improve our understanding of potential interaction effects between AI use, AI literacy and specific disclosure conditions based on the data from Table 2. Additionally, the geographical scope of our study was restricted. This limitation means that our findings are not generalizable and are therefore context-dependent. Future research would benefit from incorporating views from individuals across different media systems and various continents, ensuring a more holistic understanding of the 'disclosure puzzle'. Overall, the interviews portray AI transparency as a dynamic and

interpretive challenge rather than a fixed checklist. However, it is precisely in this complexity that opportunity lies: By engaging audiences thoughtfully and designing disclosures with their needs in mind, journalism can navigate this new technological frontier in ways that enhance - not erode - public trust.

Notes

1. <https://www.newsguardtech.com/special-reports/newsbots-ai-generated-newswebsites-proliferating/>.

Ethical Approval

This study was approved by the Ethics Review Commission of (university withheld).

Author contributions

CRediT: **Hannes Cools**: Conceptualization, Data curation, Formal analysis, Methodology, Project administration, Visualization, Writing – original draft, Writing – review & editing; **Sophie Morosoli**: Conceptualization, Data curation, Formal analysis, Investigation, Methodology, Supervision, Validation, Visualization, Writing – original draft, Writing – review & editing; **Laurens Naudts**: Conceptualization, Formal analysis, Investigation, Methodology, Validation, Writing – original draft, Writing – review & editing; **Karthikeya Venkatraj**: Conceptualization, Formal analysis, Methodology, Supervision, Validation, Writing – original draft, Writing – review & editing; **Claes de Vreese**: Conceptualization, Writing – original draft, Writing – review & editing; **Natali Helberger**: Conceptualization, Writing – original draft, Writing – review & editing.

Disclosure Statement

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

Funding

No funding was received.

ORCID

Hannes Cools  <http://orcid.org/0000-0002-4626-1527>

References

- Ahva, L., H. Heikkilä, and R. Kunelius. 2015. "Civic Participation and the Vocabularies for Democratic Journalism." In *The Routledge Companion to Alternative and Community Media* (pp. 155–164). London: Routledge.
- Altay, S., and F. Gilardi. 2024. "People Are Skeptical of Headlines Labeled as AI-Generated, Even If True or Human-Made, Because They Assume Full AI Automation." *PNAS Nexus* 3 (10): 403. <https://doi.org/10.1093/pnasnexus/pgae403>

- Ananny, M., and K. Crawford. 2018. "Seeing without Knowing: Limitations of the Transparency Ideal and Its Application to Algorithmic Accountability." *New Media & Society* 20 (3): 973–989. <https://doi.org/10.1177/1461444816676645>
- Becker, K. B., F. M. Simon, and C. Crum. 2024. "Policies in Parallel? A Comparative Study of Journalistic AI Policies in 52 Global News Organisations." *Digital Journalism* 1 (1): 1–21.
- Beckett, C. 2019. "New Powers, New Responsibilities. A Global Survey of Journalism and Artificial Intelligence." <https://blogs.lse.ac.uk/polis/2019/11/18/new-powers-new-responsibilities/>
- Beckett, C., and M. Yaseen. 2023. "Generating Change. A Global Survey of What News Organisations Are Doing with AI." <https://www.aiunplugged.io/wp-content/uploads/2023/10/Generating-Change-A-global-survey-of-what-news-organisations-are-doing-with-AI-By-Cyber-Gear.pdf>
- Braun, V., and V. Clarke. 2006. "Using Thematic Analysis in Psychology." *Qualitative Research in Psychology* 3 (2): 77–101. <https://doi.org/10.1191/1478088706qp0630a>
- Campbell, J. L., C. Quincy, J. Osseman, and O. K. Pedersen. 2013. "Coding in-Depth Semistructured Interviews: Problems of Unitization and Intercoder Reliability and Agreement." *Sociological Methods & Research* 42 (3): 294–320. <https://doi.org/10.1177/0049124113500475>
- Caswell, D. 2023. "AI and News: What's Next?." Available at: <https://generative-ai-newsroom.com/ai-and-news-whats-next-154fb6b6a646>
- Commissariaat voor de Media. 2024. "Media Monitor 2024." Available at: <https://www.cvdm.nl/wp-content/uploads/2025/02/Summary-Media-Monitor-2024.pdf>
- Cools, H., and N. Diakopoulos. 2023. "Towards Guidelines for Guidelines on the Use of Generative AI in Newsrooms." Available at: <https://generative-ai-newsroom.com/towards-guidelines-for-guidelines-on-the-use-of-generative-ai-in-newsrooms-55b0c2c1d960>
- Cools, H., and N. Diakopoulos. 2026. "Uses of Generative AI in the Newsroom: Mapping Journalists' Perceptions of Perils and Possibilities." *Journalism Practice* 20 (3): 878–896. <https://doi.org/10.1080/17512786.2024.2394558>
- de-Lima-Santos, M. F., W. N. Yeung, and T. Dodds. 2025. "Guiding the Way: A Comprehensive Examination of AI Guidelines in Global Media." *AI & SOCIETY* 40 (4): 2585–2603. <https://doi.org/10.1007/s00146-024-01973-5>
- Deuze, M. 2005. "What is Journalism?" *Journalism* 6 (4): 442–464. <https://doi.org/10.1177/1464884905056815>
- D'haeseleer, S., K. Van Damme, H. Cools, S. Van Leuven, and T. Evens. 2025. "AI Divides in Newsrooms? How Journalists in the Low Countries Use and Perceive Generative AI." *Journalism Practice* 2025: 1–28. <https://doi.org/10.1080/17512786.2025.2538120>
- Diakopoulos, N. 2015. "Algorithmic Accountability: Journalistic Investigation of Computational Power Structures." *Digital Journalism* 3 (3): 398–415. <https://doi.org/10.1080/21670811.2014.976411>
- Dobber, T., S. Kruijemeier, F. Votta, N. Helberger, and E. P. Goodman. 2025. "The Effect of Traffic Light Veracity Labels on Perceptions of Political Advertising Source and Message Credibility on Social Media." *Journal of Information Technology & Politics* 22 (1): 82–97. <https://doi.org/10.1080/19331681.2023.2224316>
- El Ali, A., K. P. Venkatraj, S. Morosoli, L. Naudts, N. Helberger, and P. Cesar. 2024, May. Transparent AI disclosure obligations Whowhatwhen, where, why, how. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*, pp. 1–11.
- Fletcher, R., and R. Nielsen. 2024. "What Does the Public in Six Countries Think of Generative AI in News?." Available at: <https://reutersinstitute.politics.ox.ac.uk/what-does-public-six-countries-think-generative-ai-news>
- Gamage, D., D. Sewwandi, M. Zhang, and A. K. Bandara. 2025. "Labeling Synthetic Content: User Perceptions of Label Designs for AI-Generated Content on Social Media." In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems* (pp. 1–29).
- Graefe, A., M. Haim, B. Haarmann, and H. B. Brosius. 2018. "Readers' Perception of Computer-Generated News: Credibility, Expertise, and Readability." *Journalism* 19 (5): 595–610. <https://doi.org/10.1177/1464884916641269>
- Heim, K. 2024. "Transparency in Digital Journalism." In *The Routledge Companion to Digital Journalism Studies*, 60–69. Routledge.
- Jones, B., E. Luger, and R. Jones. 2023. "Generative AI & Journalism: A Rapid Risk-Based Review."

- Jung, M. A. 2026. "Transparent but Not Always Trusted: The Impact of AI Disclosure on Ad Credibility." Master's Dissertation. Universidade Católica Portuguesa. Veritati Repository. <http://hdl.handle.net/10400.14/57835>
- Karlsson, M. 2010. "Rituals of transparency: Evaluating online news outlets' uses of transparency rituals in the United States, United Kingdom and Sweden." *Journalism Studies* 11 (4): 535–545. [10.1080/14616701003638400](https://doi.org/10.1080/14616701003638400).
- Karlsson, M. 2020. "Dispersing the Opacity of Transparency in Journalism on the Appeal of Different Forms of Transparency to the General Public." *Journalism Studies* 21 (13): 1795–1814. <https://doi.org/10.1080/1461670X.2020.1790028>
- Kitzinger, J. 1995. "Qualitative Research: Introducing Focus Groups." *BMJ (Clinical Research ed.)* 311 (7000): 299–302. <https://doi.org/10.1136/bmj.311.7000.299>
- Koliska, M. 2022. "Trust and Journalistic Transparency Online." *Journalism Studies* 23 (12): 1488–1509. <https://doi.org/10.1080/1461670X.2022.2102532>
- Longoni, C., A. Fradkin, L. Cian, and G. Pennycook. 2022. June). "News from Generative Artificial Intelligence is Believed Less. In *Proceedings*." of the 2022 ACM Conference on Fairness, Accountability, and Transparency (pp. 97–106). <https://doi.org/10.1145/3531146.3533077>
- Mattis, N., K. Kieslich, and C. de Vreese. 2025. Feeling iffy about generative AI When journalists disclose AI use trust in news is lower. *OSF Preprints*, 1–26.
- Morgan, D. L. 1997. *Focus Groups as Qualitative Research* (Vol. 16). London: Sage Publishers.
- Morosoli, S., Naudts, L., Cools, H., Venkatraj, K., Helberger, N., & de Vreese, C. 2026. "Public accountability and regulatory expectations for AI in journalism: qualitative evidence from focus groups with Dutch citizens." *AI & SOCIETY*, 41 (2): 1467–1479.
- Newman, N., R. Fletcher, C. T. Robertson, A. Ross Arguedas, and R. K. Nielsen. 2024. "Reuters Institute Digital News Report 2024." *Reuters Institute for the Study of Journalism*.
- Piasecki, S., S. Morosoli, N. Helberger, and L. Naudts. 2024. "AI-Generated Journalism: Do the Transparency Provisions in the AI Act Give News Readers What They Hope for?" *Internet Policy Review* 13 (4): 1–28. <https://doi.org/10.14763/2024.4.1810>
- Plaisance, P. L. 2007. "Transparency: An Assessment of the Kantian Roots of a Key Element in Media Ethics Practice." *Journal of Mass Media Ethics* 22 (2-3): 187–207. <https://doi.org/10.1080/08900520701315855>
- Prajod, P., H. Cools, T. Röggl, P. Cesar, and A. E. Ali. 2026. "Towards Gaze-Informed AI Disclosure Interfaces: Eye-Tracking Attentional and Cognitive Load While Reading AI-Assisted News.." *arXiv preprint arXiv:2605.14999*.
- Schilke, O., and M. Reimann. 2025. "The Transparency Dilemma: How AI Disclosure Erodes Trust." *Organizational Behavior and Human Decision Processes* 188: 104405. <https://doi.org/10.1016/j.obhdp.2025.104405>
- Schnell, K. 2024. "AI Transparency in Journalism: Labels for a Hybrid Era." https://reutersinstitute.politics.ox.ac.uk/sites/default/files/2025-01/RISJ%20Fellows%20Paper_Katja%20Schell_MT24_Final.pdf
- Simon, F. 2024. "Artificial Intelligence in the News: How AI Retools, Rationalizes, and Reshapes Journalism and the Public Arena." Available at: <https://ora.ox.ac.uk/objects/uuid:aeb25013-1d17-40b2-b471-5bdca309db87>
- Singer, J. B., A. Hermida, D. Domingo, A. Heinonen, S. Paulussen, T. Quandt, Z. Reich, and M. Vujnovic. 2011. *Participatory Journalism: Guarding Open Gates at Online Newspapers*. Wiley-Blackwell. <https://doi.org/10.1002/9781444340747>
- Stenbom, A., M. Wiggberg, and T. Norlund. 2023. "Exploring Communicative AI: Reflections from a Swedish Newsroom." *Digital Journalism* 11 (9): 1622–1640. <https://doi.org/10.1080/21670811.2021.2007781>
- Strikovic, E., and H. Cools. 2025. "Reality Re-Imag (in) ed. Mapping Publics' Perceptions and Evaluations of AI-Generated Images in News Contexts." *Digital Journalism* 1 (1): 1–22. <https://doi.org/10.1080/21670811.2025.2573076>
- Swart, J., T. Groot Kormelink, I. Costera Meijer, and M. Broersma. 2022. "Advancing a Radical Audience Turn in Journalism. Fundamental Dilemmas for Journalism Studies." *Digital Journalism* 10 (1): 8–22. <https://doi.org/10.1080/21670811.2021.2024764>

- Trattner, C., S. L. Forstner, A. D. Starke, and E. Knudsen. 2026. "C2PA Provenance Labels Increase Trust in News Platforms Across Western Countries."
- Venkatraj, K. P., S. Morosoli, H. Cools, L. Naudts, C. de Vreese, N. Helberger, ... A. El Ali. 2025. December). "Understanding AI Disclosure Needs for News Production and Journalism." In *Proceedings of the 24th International Conference on Mobile and Ubiquitous Multimedia*, pp. 202–208.
- Vos, T. P., and S. Craft. 2017. "The Discursive Construction of Journalistic Transparency." *Journalism Studies* 18 (12): 1505–1522. <https://doi.org/10.1080/1461670X.2015.1135754>
- Waddell, T. F. 2026. "The How and Why of Generative AI: The Effect of AI Disclosure Policies on the Perceived Trustworthiness and Financial Value of Automated News." *Mass Communication and Society* 2026: 1–17. <https://doi.org/10.1080/15205436.2026.2679456>
- Westlund, O., and S. C. Lewis. 2014. "Agents of Media Innovations: Actors, Actants, and Audiences." *The Journal of Media Innovations* 1 (2): 10–35. <https://doi.org/10.5617/jmi.v1i2.856>
- Wölker, A., and T. E. Powell. 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>
- Zamith, R. 2019. "Transparency, Interactivity, Diversity, and Information Provenance in Everyday Data Journalism." *Digital Journalism* 7 (4): 470–489. <https://doi.org/10.1080/21670811.2018.1554409>
- Zamith, R. 2024. "Open-Source Repositories as Trust-Building Journalism Infrastructure: Examining the Use of GitHub by News Outlets to Promote Transparency, Innovation, and Collaboration." *Digital Journalism* 12 (7): 985–1006. <https://doi.org/10.1080/21670811.2023.2202873>
- Zier, J., and N. Diakopoulos. "Labeling AI-Generated News Content: Matching Journalist Intentions with Audience Expectations." Available at: https://cplusj2024.github.io/papers/CJ_2024_paper_14.pdf