

Strategic Framing of AI Implementation in Media organisations: A Vignette Study of Managerial Communication and Employee Responses

Melise Peruchini ¹, Stephan Böhm ², Julio Monteiro Teixeira ³

¹Information Systems Department, Federal University of Santa Catarina, Florianópolis, Brazil

²Center for Advanced E-Business Studies, RheinMain University of Applied Sciences and Arts, Wiesbaden, Germany

³Graphic Expression Department, Federal University of Santa Catarina, Florianópolis, Brazil

e-mail: melisepe.peruchini@ufsc.br, stephan.boehm@hs-rm.de, julio.monteiro@ufsc.br

Abstract—As Artificial Intelligence (AI) becomes increasingly integrated into the daily operations of media organisations, this transformation directly affects employees' roles, responsibilities, and professional autonomy. The aim of this study is to investigate how internal management communication might shape employee perceptions and AI adoption by examining how different strategic framings of AI influence media employees' perceived threat, perceived organisational support, and perceived sincerity/credibility and evaluation of management communication. Moreover, the study aims to analyse how media employees' perceptions of AI adoption regarding trust, perceived usefulness, and job insecurity affect the intention to use AI tools. The study employs a vignette-based survey experiment using a within-subjects design, in which respondents evaluate four scenarios (framings) of managerial communication and complete the same set of questions after each one. Data were collected through an online questionnaire administered by a panel provider, targeting media employees (n = 156), and statistically analysed using correlation, linear regression, and Friedman tests with post hoc pairwise comparisons. The results suggest that AI adoption in media organisations is driven more by trust and perceived usefulness than by job insecurity, whereas strategic framing influenced only perceived organisational support.

Keywords—media organisations; artificial intelligence; AI-based transformation; strategic framing; management communication; vignette methodology.

I. INTRODUCTION

Media companies across the world are increasingly integrating Artificial Intelligence (AI) into content production and distribution, as well as into their organisational workflows. This transformation is unfolding against a backdrop of mounting economic pressures, diminishing public trust in institutions, and the increasing fragmentation of media markets. While prior studies have concentrated mainly on audience perceptions of AI-generated media content, the influence of internal communication and governance related to the introduction of AI on employee trust, perceptions of threat, professional orientation, and willingness to utilise these technologies remains an emerging area of research, as discussed in the related studies presented in Section II - Theory.

This research gap is of significant concern for the field of media management practice, as artificial intelligence is becoming increasingly integrated into the daily operations of media organisations. AI integration at corporate level has a direct impact on the organisation's legitimacy and their long-term sustainability. To address this, the present study investigates how top management's strategic framing and communication

of AI initiatives influences media employees' perceptions and intentions to use AI tools at work.

The present study is situated at the intersection of three significant fields of enquiry: (1) artificial intelligence-driven media disruption, (2) trust and credibility in media companies, and (3) media management and governance. It explores legitimacy, motive attributions, aspects of trust in company management, threat assessments, the perceived usefulness of AI, and behavioural intention [1] to use AI tools in day-to-day work. In addition, the study considers organisational context variables, including Perceived Organisational Support (POS) [2] and job insecurity [3]. These variables are particularly salient in conditions of economic uncertainty and amid AI disruption in the media industry.

The present study contributes to media management research in three ways: Firstly, it advances understanding of AI-based disruption within media organisations by redirecting the focus from audience reactions to internal management and organisational processes. Secondly, it provides an analysis of the dynamics of trust and legitimacy by showing how the design of an AI strategy can either support or undermine employees' willingness to engage with AI technologies. Thirdly, it provides empirical evidence on the effects of different strategic framings of AI implementation on employees' perceptions, contributing to the ongoing discussion about the role of managerial communication in AI adoption.

Following this introduction, Section II presents the theoretical foundations and related work. Section III describes the methodology and research design, while Section IV presents results and discussions. Finally, Section V presents the conclusions, limitations, and directions for future research.

II. THEORY

The theoretical framework of the present study draws on three strands of literature that examine media organisations in the context of technological change. Firstly, with reference to sensemaking and sensegiving theory [4], the management communication about the introduction of AI is conceptualised as a strategic attempt to influence employees' interpretations of technological disruption. Previous and related studies have focused on analysing framing effects and how they positively or negatively affect AI adoption across different fields [5] [6] [7] but comparatively little attention has been paid to media management in a broader sense. Secondly, the present

study uses organisational legitimacy theory [8], attribution theory [9], and research on reciprocal trust dynamics [10] to explain the reasons why different AI strategies can result in divergent assessments of the motives, appropriateness, and acceptance of AI integration. Thirdly, the approach incorporates organisational trust-building theory, operationalised through behavioural intentions, i.e., the willingness to be vulnerable or entrust [1], and identity threat theory [11] to capture professional concerns related to automation, rationalisation, and the perceived erosion of editorial or professional autonomy. Trust is considered a key component in technology adoption, grounded in well established models, such as Technology Acceptance Model (TAM) [12] and Unified Theory of Acceptance and Use of Technology (UTAUT) [13]. Moreover, according to [14], trust has become a crucial factor in the perception, use, and adoption of AI, having a significant positive impact on the continued use of AI and the perception of its usefulness.

Prior studies on audience perceptions of AI-generated content have mainly focused on questions of trust, authenticity, credibility, and perceived objectivity. According to [15], people sometimes read AI-generated media as more automated, colder, and more error-prone, which undermines trust; but they also sometimes see it as more objective and less tied to human political or economic pressures. AI-generated marketing content lowers perceived trust, authenticity, and credibility when consumers recognise it as AI-made [16]; and adoption depends more on perceived value and social influence than on credibility alone, while transparency, quality, and human supervision can soften rejection [17]. However, despite communication having been identified as an important factor in AI adoption within organisations [18] there is a gap in exploring AI within internal communication [19], especially considering the ongoing technological changes.

To analyse the relationships between these organisational dynamics and technology-related outcomes, the study also draws on established research on technology acceptance [12] [13]. In this study, trust refers primarily to confidence in management's responsible implementation and governance of AI whereas perceived usefulness captures the instrumental value of AI for work performance. POS, in contrast, reflects employees' perceptions that the organization considers their well-being during the AI implementation process.

Therefore, based on prior research on technology adoption and organisational communication, the following hypotheses are proposed:

H1: Trust is associated with employees' Behavioural Intention to use AI tools at work.

H2: Perceived Usefulness is associated with employees' Behavioural Intention to use AI tools at work.

H3: Strategic framing of AI implementation influences employees' perceptions of managerial communication.

The core of this study is a vignette methodology approach, which involves presenting respondents with realistic, systematically varied scenarios to examine how specific factors influence their perceptions, attitudes and Behavioural Intentions

[20]. H1 and H2 are examined through correlation and linear regression analysis based on the pre-vignette questions. H3 is tested using non-parametric Friedman tests to compare employees' evaluations across the different vignette conditions in a within-subjects design.

III. METHODS

Methodologically, this study uses an experimental vignette approach to construct hypothetical scenarios presented to respondents. According to [20], a quantitative vignette study consists of a vignette experiment as the core element, and a traditional survey for additional co-variables in the vignette data, for example, respondent specific characteristics. The study follows a within-subjects design, in which each respondent is exposed to all four vignettes (framings) developed for this research. This design allows for the comparison of individual responses across different AI strategy framings while controlling for between-subject variability [20].

For the data collection, an online questionnaire was designed, administered by a panel provider, targeting media employees through different sectors, such as journalism, broadcasting media, online media, etc. To ensure panel quality, the questionnaire included screening filters for media professionals with AI-related work experience, an attention-check item, and manipulation/comprehension checks. Incomplete or low-quality responses were removed during data cleaning. The survey was conducted in Germany, Austria, Switzerland, and the United Kingdom and administered in English and German between 16 April and 27 April 2026.

The experimental manipulation consists of four vignettes presenting internal management communications about the implementation of generative AI in a media company, all of them sharing a common introductory context and a different message framing, emphasising: (1) cost efficiency, (2) quality and verification, (3) audience engagement and responsiveness, and (4) mission and responsibility. The selection of these framings is grounded in organisational communication and legitimacy literature and reflects distinct strategic narratives commonly used in technology adoption contexts, namely efficiency-oriented, quality-oriented, audience-oriented, and responsibility-oriented arguments [8] [13].

In the vignette introduction section of the questionnaire, participants were asked to imagine that their media organisation was introducing a secure environment for the use of general-purpose AI tools (e.g., similar to ChatGPT) in compliance with data protection regulations. The tools were described as supporting tasks, such as drafting, summarizing, translating, idea generation, social media content creation, and initial analyses across both creative and non-creative roles. Participants were subsequently presented with the four vignette reflecting different strategic framings of the AI implementation and were asked to evaluate each message:

Cost efficiency: *In recent months, we have carefully reviewed how work processes are organized across teams and where time and effort is currently invested. Many activities still require substantial manual input, even when they follow*

recurring patterns. With the introduction of generative AI, we want to streamline these processes and reduce the amount of repetitive work in everyday tasks. In your daily work, this may involve using AI to prepare first drafts, structure information, or support routine production steps. We expect this to make workflows more consistent and allow teams to handle increasing demands more effectively, while creating more room for tasks requiring judgment and expertise.

Quality & verification: The increasing speed of content production has made it more challenging to maintain consistent standards across all outputs. Even small inaccuracies or inconsistencies can affect how our work is perceived. With the introduction of generative AI, we aim to support you in maintaining a high level of accuracy and coherence. In your daily work, AI tools can help structure content, identify inconsistencies, and provide additional reference points when working with information. We see this as a way to strengthen the overall reliability of our work and provide additional support where careful review is particularly important.

Audience engagement: Patterns of content use and interaction have become more dynamic, and it is increasingly important to recognize how different formats perform in practice. With the introduction of generative AI, we want to strengthen our ability to work with these patterns in a more flexible and responsive way. In your daily work, this may involve adapting content more quickly, experimenting with alternative formats, and using available data more systematically to inform decisions. This should make it easier to adjust ongoing work based on observed responses and to refine outputs continuously in a fast-changing environment.

Mission & responsibility: As a media organization, our work is closely connected to broader expectations regarding responsibility, credibility, and transparency. These expectations continue to evolve as new technologies are introduced. With the introduction of generative AI, we want to ensure that its use is aligned with our professional standards and guiding principles. In your daily work, this includes following clear guidelines, documenting the use of AI where relevant, and maintaining oversight over outputs. Our aim is to integrate these tools in a way that supports consistent practices and reinforces expectations associated with our role.

After each vignette, respondents rate on a 5-point Likert scale items including manipulation and comprehension / realism checks (to assess whether the intended differences between the vignettes were perceived and to ensure comprehension and response validity), followed by measures regarding the organisation's message, such as Evaluation of the Communication, Perceived Sincerity / Credibility, Perceived Threat, and POS. The repeated measures are collected to capture communication-specific reactions through the within-subject design, enabling comparisons within individuals, assessing how different communication approaches influence their perceptions. The order of vignette presentation was randomized across participants.

As mentioned, the proposed questionnaire also aims to consider aspects of Trust, Behavioural Intention, Perceived

Usefulness, and Job Insecurity; therefore, general questions about the use of AI tools in daily work were included. These constructs are well established in the technology adoption literature, and the measures were adapted from prior research, including the Behavioural Trust Inventory (BTI) [1] TAM for Perceived Usefulness [12], and UTAUT for Behavioural Intention [13], as well as prior research on Job Insecurity [21]. All constructs were measured using multi-item scales, mostly using two items, while trust was measured using four items for capturing dimensions of reliance and disclosure. The separation between general and framing-specific measures was implemented to reduce respondent fatigue and minimize redundancy and achieve a clearer distinction between stable attitudes toward AI and context-dependent evaluations. Table I represents each construct measured using the questionnaire.

TABLE I. CONSTRUCTS AND QUESTIONS

Construct	Example Item
<i>Pre-framing measures (general attitude studies)</i>	
Behavioural intention	I would be willing to use AI tools in my work.
Trust (reliance and disclosure)	I would feel comfortable sharing work-related information when using AI tools; I believe management will use AI responsibly.
Perceived usefulness	The use of AI would increase organizational effectiveness.
Job insecurity	I am concerned about my job security due to the use of AI.
<i>Framing-based measures (repeated for each framing)</i>	
Manipulation check	The message mainly emphasizes efficiency considerations
Comprehension check	The message reflects a realistic situation in a media organization
Evaluation of communication	"The management message is convincing
Perceived Sincerity/Credibility	The reasons presented by management appear to reflect the actual motives behind this initiative
Perceived Threat	This message increases my uncertainty regarding my future role
Perceived Organizational Support	The management message considers employees' well-being

Following the operationalisation of constructs presented above, the data were analysed using R and the RStudio environment. Prior to analysis, the dataset was cleaned to ensure completeness, response quality, and the removal of potential outliers, resulting in a final sample of 156 respondents. Subsequently, the constructs were assessed for internal consistency using Cronbach's alpha. All scales demonstrated acceptable to excellent reliability, supporting the aggregation of items into composite scores. Following this, correlation and linear regression analyses were conducted to examine the relationships between Trust, Perceived Usefulness, Job

Insecurity, and Behavioural Intention. For the framing-based measures, non-parametric Friedman tests were applied to assess differences across the four framings. The results of these analyses are presented and discussed in detail in Section IV, where the implications of the findings are interpreted in relation to the research objectives and hypothesis.

IV. RESULTS AND DISCUSSION

This section is divided into three subsections: descriptive statistics, pre-framing measures, and framing effects.

A. Descriptive Statistics

Regarding the demographic characteristics, the sample is relatively balanced in terms of gender (male 46%, female 53%, diverse 1%), while age is mainly concentrated in ranges approximately 25–44 years. A substantial proportion of participants report holding higher education degrees, e.g., on bachelor's (21%) and master's levels (33%). Respondents have extensive experience in the media industry, (60% more than 10y) and are mostly employed within organisations (77%), self-employed individuals (12%) and freelancers (10%). Positions are primarily at the employee level (39%) and middle management (25%), with smaller shares in top management (11%). In terms of company size, most respondents work in smaller organisations with up to 99 employees (46%), and the predominant areas are print and internet-based media (above 50%). General attitude towards the use of AI in work processes were relatively positive (48.41%), followed by neutral views (34.92%) and negative attitudes (16.67%).

When asked about the potential for task substitution by AI in the present, participants responses were widely dispersed, concentrating around lower to moderate levels (e.g., 20% and 40%). Expectations for long-term substitution were generally higher (values as 40% and 80%) than perceptions of current replaceability, as shown in Figure 1.

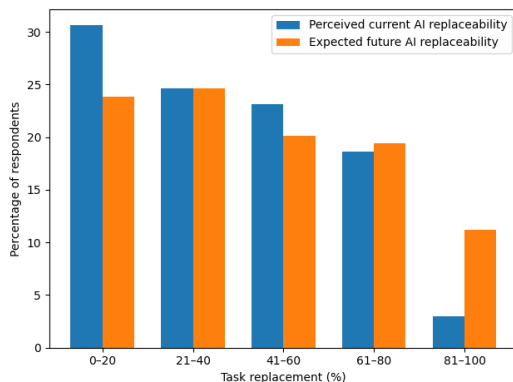


Figure 1. Perceived and expected AI replaceability

The figure illustrates the perceptions of the respondents about their own tasks that can currently be replaced by AI (“For what percentage of your own tasks and activities do you already consider replacement by AI to be fundamentally possible today?”) and their expectations regarding long-term replacement (“For what percentage of your own tasks do you expect long-term replacement by AI?”). This suggests that

they perceive the current potential for AI replaceability as moderate, while expecting AI to affect a substantially larger share of tasks in the future.

B. Pre-framing Measures

As mentioned in the methods section, prior to the vignette experiment, respondents were presented with a general framing and asked a set of questions capturing four key constructs: Behavioural Intention, Trust, Perceived Usefulness, and Job Insecurity. Based on this scenario, respondents evaluated their behavioural associated with the implementation of AI. The internal consistency of the multi-item constructs was assessed using Cronbach’s alpha before aggregation. All scales demonstrated acceptable to excellent reliability, as presented in table II:

TABLE II. DESCRIPTIVE STATISTICS AND RELIABILITY MEASURES

Construct	M	SD	α
Behavioral intention	1.90	0.84	.88
Trust	2.20	0.75	.81
Perceived usefulness	2.40	0.85	.77
Job insecurity	3.00	1.20	.93

Given these results, all items were retained and aggregated by computing mean scores for each construct.

Correlation analysis revealed that Behavioural Intention to use AI is moderately and positively associated with trust ($r = .51$) and perceived usefulness ($r = .58$). In contrast, job insecurity exhibits only a weak negative relationship with behavioural intention ($r = -.10$).

To further examine these relationships, a linear regression analysis was conducted with Behavioural Intention as the dependent variable. The regression diagnostics indicated low values of variance inflation factor (all VIFs < 1.40) and the residual plots showed no substantial deviations from homoscedasticity or normality. The regression results indicate that both Trust ($\beta = .31$, $t(152) = 3.78$, $p < .001$) and Perceived Usefulness ($\beta = .43$, $t(152) = 5.97$, $p < .001$) are significant predictors of Behavioural Intention. Job Insecurity, however, does not show a significant effect ($\beta = -.04$, $t(152) = -0.81$, $p = .42$). The analysis revealed a statistically significant regression effect, $F(3, 152) = 34.08$, $p < .001$, accounting for 40% of the variance in Behavioural Intention ($R^2 = .40$, adjusted $R^2 = .39$).

Therefore, the data suggest that, in this sample, employees intend to use AI when they trust it or perceive it as useful, while perceptions of potential future impacts on their jobs do not directly affect their intention to use it. In summary, AI adoption appears to be driven more by positive perceptions than by job insecurity.

C. Framing Effect

Following the analysis of the general and pre-framing constructs, the responses to the framing-based measures were examined. The manipulation checks revealed statistically significant differences between framing conditions ($\chi^2 = 8.38$,

$df = 3$, $p = .039$), indicating that participants were able to distinguish between them. For the comprehension (I clearly understood the message from management) and realism (The message reflects a realistic situation in a media organisation) check, participants generally reported that the framing messages were understandable across all conditions: Between 71% and 84% of respondents agreed or strongly agreed that they clearly understood the management messages, with mean values ranging from $M = 1.97$ to $M = 2.19$. Realism was slightly lower but remained relatively positive across conditions, with 65% to 73% and $M = 2.19$ – 2.33 .

For the remaining framing-based measures, Cronbach's alpha shows acceptable to excellent reliability: Evaluation of the Communication ($\alpha = .77$ – $.86$), Perceived Sincerity/Credibility ($\alpha = .74$ – $.85$), Perceived Threat ($\alpha = .78$ – $.93$), and POS ($\alpha = .78$ – $.85$).

Given the within-subject design and the ordinal nature of the Likert-scale data, non-parametric Friedman tests were conducted to assess whether the different framings led to statistically significant differences in respondents' evaluations. The analysis included the constructs Evaluation of Communication, Perceived Sincerity/Credibility, Perceived Threat, and POS.

Across the four constructs, the results did not reveal statistically significant differences between framing conditions for Evaluation of Communication ($\chi^2 = 3.36$, $p = .34$), Perceived sincerity/credibility ($\chi^2 = 2.27$, $p = .52$), or Perceived Threat ($\chi^2 = 3.12$, $p = .37$). However, POS showed statistically significant differences across conditions ($\chi^2 = 8.31$, $p = .040$), therefore, the different framing significantly influenced POS, indicating that employees interpreted organisational support differently depending on how the AI implementation was framed. The results of the Friedman tests are presented in Table III.

TABLE III. FRIEDMAN TEST RESULTS ACROSS FRAMING CONDITIONS

Construct	χ^2	df	p-value
Evaluation of the Communication	3.36	3	.339
Perceived Sincerity and Credibility	2.27	3	.518
Perceived Threat	3.12	3	.373
Perceived Organisational Support	8.31	3	.040*

The Post-hoc pairwise comparisons revealed that the significant difference in POS was specifically observed between the quality and verification framing and the mission and responsibility framing ($p = .027$). Participants perceived higher levels of organisational support when AI implementation was framed in terms of improving accuracy, consistency, and verification processes than when it was framed around mission and responsibility of the organisation. One possible explanation for the limited framing effects is that employees may already hold relatively stable attitudes toward AI, making them less sensitive to variations in managerial communication. It is important to highlight that participants were able to distinguish between the framings, however, organizational culture and prior experiences with AI may play a stronger role in shaping

perceptions than strategic messaging alone. The findings may therefore suggest that concrete organizational implementation practices are more influential than communication framing in influencing employees' attitudes toward AI.

Given that statistically significant differences between framing conditions were limited to POS, the analysis was extended using a descriptive approach aggregating responses across all of them. This step provides, therefore, an overall description of distribution and central tendencies of responses across the key constructs.

Figure 2 illustrates the aggregated response patterns for each construct, with all items measured using a five-point Likert scale ranging from 1 ("strongly agree") to 5 ("strongly disagree"). The figure

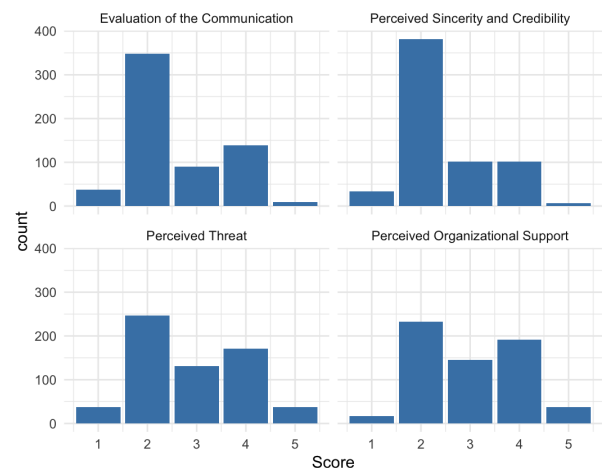


Figure 2. Frequencies aggregated across the four vignettes (4 × N=156)

The descriptive results indicate that the Evaluation of Communication and Perceived Sincerity/Credibility were generally positive, with most responses concentrated in the agreement categories of the scale. The findings represented in Figure 2 suggest that respondents tended to perceive the management messages as clear, understandable, reasonably sincere and credible. Perceived Threat and POS showed a more dispersed distribution, although many respondents selected agreement-oriented responses, a considerable proportion also chose neutral or disagreement categories, showing a greater variation in how participants perceived the potential risks and implications of AI implementation for employees and their professional roles. POS presented a moderate tendency toward agreement, however, compared to the other constructs, responses were less concentrated, indicating less consensus among participants regarding the extent to which the organisation appeared supportive during the AI implementation process.

V. CONCLUSION AND FUTURE WORK

In conclusion, the findings indicate that media employees' intention to use AI is primarily driven by trust and perceived usefulness, while job insecurity does not play a significant role within this context. Moreover, different strategic framings of AI implementation do not significantly influence employees'

Evaluations of Communication, Perceived Threat or Perceived Sincerity/Credibility. However, AI framing can influence POS. Overall, in this study, AI adoption appears to be shaped more by positive perceptions than by concerns about potential risks. Therefore, this experiment suggests that the effectiveness of AI implementation in media organisations may depend less on how strategies are framed and more on whether organisations succeed in fostering trust and demonstrating tangible usefulness in everyday work practices. The lack of significant framing effects may also point to a potential limitation in how employees perceive or differentiate managerial narratives, indicating that communication alone may not be sufficient to shape attitudes toward AI. From a practical perspective, the findings suggest that media organizations should prioritize trust-building and demonstrate the concrete usefulness of AI tools in everyday work practices.

These conclusions, however, should be interpreted taking into account the study's limitations. First, the sample size and focus on media professionals may restrict the generalizability of the findings to other industries. Second, the use of self-reported measures may introduce common method bias and limit causal inference. Finally, the survey enabled participation across different countries through English and German versions, therefore cultural and linguistic differences may have influenced respondents' interpretations. Future research should explore alternative mechanisms, such as participatory implementation processes or training initiatives, that may strengthen employees' engagement with AI. Additionally, refining framing manipulations or employing mixed-method approaches could help capture more nuanced perceptions and better isolate framing effects. Lastly, future research could further investigate alternative drivers of AI adoption beyond communication framing.

REFERENCES

- [1] N. Gillespie, *Measuring Trust in Work Relationships: The Behavioral Trust Inventory*. Melbourne Business School, 2003.
- [2] R. Eisenberger, R. Huntington, S. Hutchison, and D. Sowa, "Perceived organizational support," *Journal of Applied Psychology*, vol. 71, no. 3, pp. 500–507, 1986. DOI: 10.1037/0021-9010.71.3.500.
- [3] L. Greenhalgh and Z. Rosenblatt, "Job insecurity: Toward conceptual clarity," *The Academy of Management Review*, vol. 9, no. 3, pp. 438–448, 1984, ISSN: 03637425. [Online]. Available: <http://www.jstor.org/stable/258284> (visited on 05/09/2026).
- [4] D. A. Gioia and K. Chittipeddi, "Sensemaking and sensegiving in strategic change initiation," *Strategic Management Journal*, vol. 12, no. 6, pp. 433–448, 1991. DOI: 10.1002/smj.4250120604.
- [5] S. Ranjbar, "Perceiving ai in the newsroom: Gender, experience, and framing effects among iranian journalists," *SSRN Electronic Journal*, 2025. DOI: 10.2139/ssrn.5432174.
- [6] T. Kim and H. Song, "The effect of message framing and timing on the acceptance of artificial intelligence's suggestion," in *Extended Abstracts of the 2020 CHI Conference on Human Factors in Computing Systems*, ser. CHI EA '20, Honolulu, HI, USA: Association for Computing Machinery, 2020, pp. 1–8. DOI: 10.1145/3334480.3383038.
- [7] A. Xie, Q. Zhang, Z. Yu, and L. Yi, "Catalyzing ai-driven change: How perceived leaders' communication framing motivates team ai adoption," *Journal of Organizational Change Management*, pp. 1–27, Jan. 2026. DOI: 10.1108/JOCM-08-2025-0718.
- [8] M. C. Suchman, "Managing legitimacy: Strategic and institutional approaches," *The Academy of Management Review*, vol. 20, no. 3, pp. 571–610, 1995, ISSN: 03637425. [Online]. Available: <http://www.jstor.org/stable/258788> (visited on 05/07/2026).
- [9] P. Harvey, K. Madison, M. Mmartinko, T. R. Crook, and T. A. Crook, "Attribution theory in the organizational sciences: The road traveled and the path ahead," *Academy of Management Perspectives*, vol. 28, no. 2, pp. 128–146, 2014, ISSN: 15589080, 19434529. [Online]. Available: <http://www.jstor.org/stable/43822046> (visited on 05/07/2026).
- [10] D. L. Ferrin, M. C. Bligh, and J. C. Kohles, "It takes two to tango: An interdependence analysis of the spiraling of perceived trustworthiness and cooperation in interpersonal and intergroup relationships," *Organizational Behavior and Human Decision Processes*, vol. 107, no. 2, pp. 161–178, 2008. DOI: 10.1016/j.obhdp.2008.02.012.
- [11] J. L. Petriglieri, "Under threat: Responses to and the consequences of threats to individuals' identities," *Academy of Management Review*, vol. 36, no. 4, pp. 641–662, 2011.
- [12] F. D. Davis, "Perceived usefulness, perceived ease of use, and user acceptance of information technology," *MIS Quarterly*, vol. 13, no. 3, pp. 319–340, 1989, ISSN: 02767783, 21629730. [Online]. Available: <http://www.jstor.org/stable/249008> (visited on 05/07/2026).
- [13] V. Venkatesh, M. G. Morris, G. B. Davis, and F. D. Davis, "User acceptance of information technology: Toward a unified view," *MIS Quarterly*, vol. 27, no. 3, pp. 425–478, 2003. DOI: 10.2307/30036540.
- [14] F. Orbán and Á. Stefkovics, "Trust in artificial intelligence: A survey experiment to assess trust in algorithmic decision-making," *AI & Society*, vol. 40, pp. 4955–4969, 2025. DOI: 10.1007/s00146-025-02237-6.
- [15] E. Yeste-Piquer, J. Suau-Martínez, M. Sintes-Olivella, and E. Xicoy-Comas, "What if i prefer robot journalists? trust and objectivity in the ai news ecosystem," *Journalism and Media*, vol. 6, no. 2, 2025. DOI: 10.3390/journalmedia6020051.
- [16] H. Cao, "A study on the influencing factors of aigc-generated news content adoption based on audience perception," in *Proceedings of the 2025 International Symposium on Machine Learning and Social Computing*, ser. MLSC '25, New York, NY, USA: Association for Computing Machinery, 2025, pp. 14–18. DOI: 10.1145/3778450.3778453.
- [17] M. Kubovics, "Influencing the perception and acceptance of ai-generated marketing content as part of sg," *Journal of Lifestyle and SDGs Review*, vol. 5, no. 6, e06929, 2025. DOI: 10.47172/2965-730X.SDGsReview.v5.n06.pe06929.
- [18] S. Kelley, "Employee perceptions of the effective adoption of ai principles," *Journal of Business Ethics*, vol. 178, pp. 871–893, 2022. DOI: 10.1007/s10551-022-05051-y.
- [19] C. A. Yue, L. R. Men, R. Mitson, D. Z. Davis, and A. Zhou, "Artificial intelligence for internal communication: Strategies, challenges, and implications," *Public Relations Review*, vol. 50, no. 5, p. 102515, 2024, ISSN: 0363-8111. DOI: <https://doi.org/10.1016/j.pubrev.2024.102515>.
- [20] C. Atzmüller and P. M. Steiner, "Experimental vignette studies in survey research," *Methodology: European Journal of Research Methods for the Behavioral and Social Sciences*, vol. 6, no. 3, pp. 128–138, 2010. DOI: 10.1027/1614-2241/a000014.
- [21] M. Sverke, J. Hellgren, and K. Näswall, "Job insecurity: A literature review," SALTSA, Stockholm, Tech. Rep., 2006.