

# Human-Centred Explainable Conversational AI: Explainability, Persona Adaptation, and Entrepreneurial Agency

Kudzai Sauka<sup>1,2,\*</sup>

<sup>1</sup>Amsterdam University of Applied Sciences (AUAS), Faculty of Economics and Business, Amsterdam, the Netherlands

<sup>2</sup>University of Amsterdam, Faculty of Economics and Business, Section Entrepreneurship and Innovation, Amsterdam, the Netherlands

## Abstract

AI systems are increasingly embedded in decision-support contexts, where the way they explain and communicate outputs shapes whether recommendations are verified or accepted without verification. This concern intensifies as systems shift from reactive responses to agentic assistance, with suggestions and actions increasingly initiated by the system rather than the user. This creates tension around agency because it remains unclear whether these design features support regulated judgment or instead undermine adaptive decision-making in entrepreneurial contexts. This study examines whether attitudinal trust and entrepreneurial agency are dissociable in entrepreneurs' interactions with an AI decision-support assistant. It combines a human-centred track on attitudinal miscalibration, an interaction-level track on DST-inspired intent disambiguation and tone calibration, and an evidence-level track on graph-conditioned rationale artefacts for GraphRAG reranking. The central question is whether explanation and conversational design features improve calibrated reliance and preserve entrepreneurial agency, rather than merely making systems feel more trustworthy.

## Keywords

explainable AI, conversational agents, persona adaptation, entrepreneurial agency, trust calibration, retrieval-augmented generation

## 1. Context and Motivation

As Artificial Intelligence (AI) systems become part of entrepreneurial decision support, they do more than provide information: they recommend actions, frame evidence, and shape how choices are evaluated. Explanation and communication design therefore matter, as they can support regulated judgment or encourage trust that is not grounded in system reliability [1, 2, 3]. Recent work shows that entrepreneurs accept AI for routine tasks but become more sceptical when outputs cannot be examined or when decisions are consequential [1, 4]. When the quality of AI recommendations is not directly observable, calibrated reliance becomes difficult to sustain [5, 6].

The conceptual gap motivating this research emerged from the first phase of the PhD. Human-centred studies published in the Journal of Consumer Marketing (JCM, 2025) [7] and at XAI 2025 [8] show that anthropomorphic cues increase trust and adoption intent through social presence, while AI disclosure weakens this relational pathway, and that explainability strengthens cognition-based trust independently, with both cues producing stronger combined effects under higher perceived risk. The pattern is consistent with a dual-route process in which explainability signals competence while anthropomorphic cues increase social presence. Both routes were assessed using Likert-scale measures of attitudinal trust, understood as a relatively stable belief in system reliability [2]. However, attitudinal trust does not map cleanly onto reliance behaviour. Scharowski et al. [2] show that self-reported trust is a poor standalone predictor of reliance, while Klingbeil et al. [3] show that higher attitudinal trust can still increase reliance on AI advice, including when recommendations are accepted without

*Late-breaking work, Demos and Doctoral Consortium, colocated with the 4th World Conference on eXplainable Artificial Intelligence: July 01–03, 2026, Fortaleza, Brazil*

\*Corresponding author.

✉ k.sauka@hva.nl (K. Sauka)

ORCID 0000-0002-3233-895X (K. Sauka)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

scrutiny. Perceptual measures alone therefore cannot indicate whether these design features support informed decision-making or weaken it. Entrepreneurial agency is conceptualised here as the capacity for informed and adaptive judgment in dynamic decision contexts [9, 10]. In AI-mediated decision support, this is operationalised as adaptive decision behaviour, observed through evidence inspection, recommendation acceptance, and judgment revision [11]. Verification behaviour refers specifically to evidence inspection and judgment revision, particularly where perceived trust diverges from verification behaviour, indicating unregulated reliance [2, 3, 12].

Two failure points follow from this diagnosis and can be understood as a two-stage failure model in AI-mediated decision support. The first occurs at the interaction level, before retrieval. When a system commits prematurely to an interpretation, it generates responses from a wrong premise, such that even a well-formed explanation describes the wrong output. The error is difficult to detect because the explanation appears coherent, even though the underlying interpretation has not been verified. The second failure point occurs at the evidence level, after retrieval. Even when systems provide citations, provenance does not indicate whether the retrieved evidence influenced the system's response or how it was weighted in the decision process [4, 5]. At this stage, transparency framing can sustain high perceived trust even when reliability is unclear, which means explanation artefacts must be evaluated against verification behaviour rather than perceived usefulness alone.

This study develops a mechanism-level account of attitudinal-behavioural miscalibration in AI-mediated entrepreneurial decision support by examining the conditions under which increases in perceived trust fail to translate into evidence-checking behaviour. In doing so, it responds to calls to move beyond perceptual ratings in XAI evaluation and to assess whether explanation artefacts shape how users engage with system outputs rather than how trustworthy those outputs appear [2, 13]. Appropriate reliance is defined as the ability to accept correct AI advice and reject incorrect advice [14]. Within this framing, verification behaviour is treated as a necessary, but not sufficient, condition for such reliance because it indicates whether users undertake evaluative processing of system outputs, even though such engagement does not by itself guarantee correct decisions [15, 16]. Explanation design must therefore change not only system perception but also users' willingness to verify when evidence is weak. When explanation increases perceived trust without increasing verification, the same miscalibration pattern is reproduced. The interaction-level and evidence-level components are treated as instruments for testing this criterion. The underlying claim is that explanation and persona cues shape reliance by influencing how users resolve uncertainty between their own judgement and system output, so that explanations perceived as plausible or task-relevant may still increase adherence even when users remain unable to distinguish correct from incorrect recommendations [12].

## 2. Key Related Work

This study draws together research on explainable AI, human-centred trust, and entrepreneurial agency, with particular attention to how these streams intersect in AI-mediated decision support.

Explainable AI (XAI) research distinguishes technically faithful explanation artefacts from explanations that are merely persuasive or legible [17, 18]. Post-hoc methods such as local surrogates, feature attribution, and counterfactual generation have been evaluated primarily through fidelity metrics and perceived interpretability. Their extension to conversational retrieval systems remains limited in current literature. Knowledge-graph-based approaches partially address this limitation by grounding explanations in explicit entity-relation representations and reasoning paths, enabling traceable links between model outputs and underlying domain knowledge [19]. However, these approaches address how decisions can be explained, not whether such explanations influence how users verify or revise decisions in interaction. More importantly, evaluation in this literature focuses on whether explanations are understandable or faithful, not on whether they change how users engage with decisions [20]. Hon-drich and Ruschemeier [21] identify verification behaviour, specifically whether users actually check system outputs, as central to meaningful human involvement in AI-assisted decisions, yet this criterion remains largely absent from standard XAI evaluation frameworks. An explanation can therefore appear

useful while leaving behaviour unchanged.

Human-centred AI and service research on trust, anthropomorphism, and adoption offers a complementary perspective [22]. Prior work shows that perceived explainability and anthropomorphic cues influence trust formation through distinct cognitive and affective routes, with explainability signalling competence and anthropomorphic cues increasing social presence [8]. These effects can translate into higher adoption intent even when they are not grounded in system reliability, indicating that trust may become misaligned with behaviour. Research on human–AI decision-making further shows that self-reported trust is a poor predictor of reliance and that calibrated use requires users to accept correct advice and reject incorrect advice [2, 14, 13]. Related empirical work shows that explanations can influence adherence independently of correctness, meaning that users may rely more on AI outputs without improving their ability to detect errors [15, 16, 12]. Trust and behavioural reliance therefore cannot be treated as interchangeable.

Entrepreneurship research provides the third perspective [23]. Entrepreneurial agency is commonly conceptualised as the capacity of actors to intentionally shape, evaluate, and act within uncertain and structurally constrained environments [9, 10]. This capacity is expressed through adaptive judgment, opportunity enactment, and context-sensitive decision-making [11]. However, these frameworks were largely developed before the proliferation of AI and assume that cognition, search, and decision-making reside primarily within the individual entrepreneur. Recent work challenges this assumption. AI is increasingly conceptualised as a co-agent that redistributes cognition, search, and decision-making across human and machine systems [24, 25]. Generative systems can function as collaborators in ideation and planning, effectively expanding the scope of individual entrepreneurial action [26]. At the same time, studies show that users continue to steer interaction and selectively engage with system outputs, indicating that agency is not simply displaced but reconfigured in human–AI interaction [27]. These developments suggest that entrepreneurial agency in AI-mediated contexts depends not only on expanded capabilities but also on how users evaluate and regulate their reliance on system outputs [24, 25, 27].

Taken together, these three streams converge on a shared limitation. XAI research focuses on interpretability without establishing whether explanations influence behaviour. Trust research shows that perceived trust can increase independently of reliability and decision quality. Entrepreneurship research documents the shift toward co-agency but does not specify how evaluative control is exercised within human–AI configurations. The present work addresses this intersection by examining whether explanation and persona design features align perceived trust with entrepreneurial agency expressed through verification behaviour. Belief-based clarification, tone calibration, and graph-conditioned rationales are treated as instruments for testing this alignment and are assessed for their influence on verification behaviour and appropriate reliance in AI-mediated decision contexts.

### **3. Research Questions, Hypotheses, and Objectives**

The study is organised around three research questions. The first asks how explainability and persona cues shape attitudinal trust and adoption intent across varying levels of perceived risk. The second asks how DST-inspired intent disambiguation and bounded tone calibration shape how users resolve uncertainty before retrieving evidence. The third asks how evidence-level explanation artefacts influence verification behaviour, rather than passive compliance, in AI-mediated decision tasks. Within this framing, trust is treated as an attitudinal outcome, verification behaviour as an observable behavioural outcome, and reliance as the degree to which behaviour aligns with or diverges from system correctness.

These questions are examined across three linked stages. Stage 1 establishes the attitudinal baseline. Stage 2 examines interaction-level mechanisms before retrieval. Stage 3 evaluates whether evidence-level explanation artefacts influence verification behaviour in entrepreneurial decision tasks.

## **Stage 1: Attitudinal miscalibration baseline (RQ1)**

The human-centred studies examine how explanation and anthropomorphic cues shape trust formation and adoption intent. These studies establish the attitudinal baseline against which behavioural divergence in later stages is interpreted.

**H1.** Higher perceived explainability is associated with stronger cognition-based trust, as users interpret explanations as signals of system competence.

**H2.** Anthropomorphic cues increase affect-based trust by strengthening perceived social presence in the interaction.

**H3.** Affect-based trust increases adoption intent independently of system reliability, with stronger effects under higher perceived risk.

## **Stage 2: Interaction-level mechanisms (RQ2)**

The second stage examines two mechanisms that shape users' responses to AI during interactions before evidence is retrieved. One concerns intent resolution: DST-inspired belief-based clarification maintains uncertainty across turns and triggers clarification questions rather than committing prematurely to a single interpretation. The other concerns agentic communication: trait-calibrated tone shapes whether users experience system initiative as supportive of their own judgment or as displacing it.

**H4.** Making intent uncertainty explicit through belief-based clarification reduces premature commitment to incorrect interpretations and improves downstream intent resolution accuracy in intent-sensitive tasks.

**H5.** When agentic behaviour is communicated through trait-calibrated tone, users report higher perceived personal agency and autonomy comfort than under fixed tone conditions, particularly among more extraverted and agreeable users.

## **Stage 3: Evidence-level explanation artefacts and behavioural consequences (RQ3)**

The final stage evaluates whether evidence-level explanation artefacts influence verification behaviour in AI-mediated decision tasks. These artefacts are intended to expose why documents were retrieved and ranked, and whether evidence influenced the generated response. The central concern is whether explainability supports verification behaviour, while anthropomorphic and personalised cues weaken that effect through the relational trust route documented in Stage 1.

**H6.** Explainability is positively associated with evidence inspection and judgment revision, as it increases the perceived value of checking system outputs relative to accepting recommendations without scrutiny.

**H7.** Anthropomorphic design cues weaken the positive association between explainability and verification behaviour, because they activate affect-based trust through social presence and substitute relational confidence for epistemic scrutiny, consistent with the miscalibration pattern established in Stage 1.

**H8.** Personalisation amplifies the weakening effect in H7, such that the negative influence of anthropomorphic cues on the explainability-to-verification relationship is stronger when the adaptive communication style is matched to the user profile.

## **4. Research Approach, Methods, and Rationale**

The empirical design combines vignette-based survey experiments, controlled behavioural experiments in entrepreneurial decision contexts, and technical implementation of explanation artefacts. Attitudinal

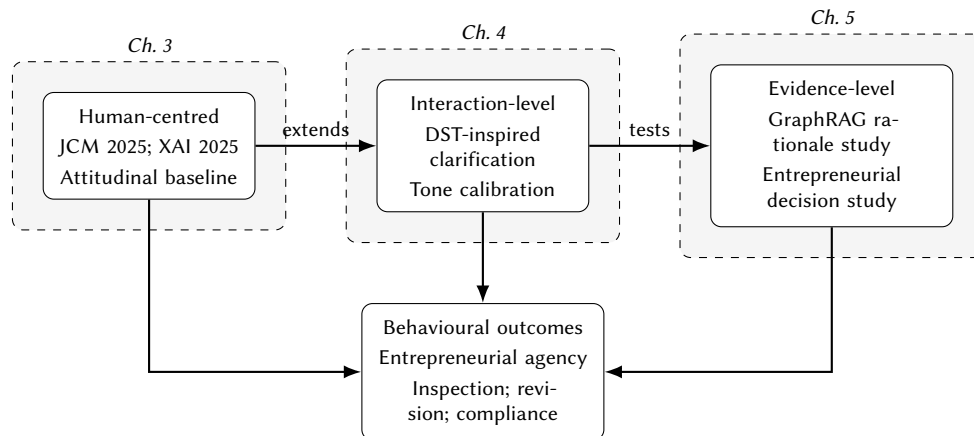
miscalibration can be identified through perceptual measures, but its consequences can only be observed by recording behaviour in decision settings. Survey methods therefore establish the baseline in Stage 1, while task-based experiments in Stages 2 and 3 allow behavioural responses to be observed directly rather than inferred from self-report alone.

The human-centred track employs  $2 \times 2$  between-subjects vignette experiments manipulating explainability and anthropomorphism in AI service encounters, with participants recruited through Prolific. Constructs are measured using established multi-item scales for perceived explainability, anthropomorphism, social presence, trust, perceived risk, and adoption intent. In addition to regression-based analysis, fsQCA is used to capture the possibility that high adoption intent may arise through different sufficient configurations of these conditions rather than through a single additive pathway.

The interaction-level track uses a controlled  $2 \times 2$  experiment manipulating agency mode (reactive vs. agentic) and tone control (fixed vs. trait-calibrated) in an entrepreneurial scenario-based task involving idea refinement and decision evaluation. Participants are recruited through Prolific ( $N = 200$ ;  $f = 0.25$ ,  $\alpha = .05$ , power = .80). The platform records accepted suggestions and delegated subtasks alongside self-reported perceptions of personal agency, autonomy, comfort, and trust. Extraversion and Agreeableness, measured using the Ten-Item Personality Inventory (TIPI), are included as targeted moderators to capture differences in how users respond to agentic behaviour, but are not treated as the central explanatory mechanism.

The evidence-level track employs a  $2 \times 2$  experiment manipulating explainability and persona adaptation in a market-entry decision scenario under uncertainty ( $N = 220$ ;  $f = 0.25$ ,  $\alpha = .05$ , power = .80). Participants evaluate competing expansion options based on incomplete and partially conflicting information. The platform logs evidence inspection, recommendation acceptance, and judgment revision. These actions are treated as behavioural proxies for reliance rather than direct measures of decision quality [16, 15]. Verification behaviour is therefore analysed as a necessary but not sufficient condition for appropriate reliance.

Across all stages, technical components such as DST-inspired clarification, tone calibration, and graph-conditioned retrieval rationales are used to implement the experimental conditions. They function as controlled instruments for identifying where miscalibration emerges across interaction and evidence levels, and are evaluated against whether they increase verification behaviour rather than unexamined acceptance of recommendations.



**Figure 1:** Three-track thesis structure. The human-centred track establishes the attitudinal miscalibration baseline (Ch. 3); the interaction-level and evidence-level tracks provide behavioural tests of whether explanation artefacts move entrepreneurs beyond that baseline (Chs. 4–5).

## 5. Preliminary Results and Contributions to Date

The human-centred track is substantially complete. The JCM 2025 study ( $n = 444$ ) shows that anthropomorphism increases adoption intent through perceived social presence and trust, while transparency framing weakens but does not remove this relational pathway. Ongoing work extending these findings with fsQCA shows that high adoption intent can emerge through either transparency-focused or relational configurations, depending on perceived risk. The XAI 2025 study [8] finds that explainability and anthropomorphism jointly increase trust, with the cognitive route remaining more stable across contexts than the relational route. Across these studies, trust can be elevated through design features independently of system reliability, thereby establishing the attitudinal miscalibration baseline.

At the interaction level, the DST-inspired intent disambiguation system presented at the ProActLLM workshop at CIKM 2025 [28] shows that maintaining competing intent hypotheses and triggering clarification under uncertainty improves intent resolution accuracy and calibration relative to softmax-based approaches. The tone calibration architecture, a related technical paper currently under review, complements this by showing that anthropomorphic tone can be decomposed into bounded parameters with controlled variation and logging. Together, these components establish that interaction-level mechanisms for managing uncertainty and social framing can be implemented and evaluated.

At the evidence level, the GraphRAG rationale study is in advanced development and focuses on graph-conditioned rationales for retrieval decisions in GraphRAG systems. Technical evaluation considers surrogate fidelity, deletion-based effects, and rank sensitivity to assess whether the rationales reflect the underlying retrieval process. The behavioural component centres on a market-entry decision scenario in which participants evaluate competing options under uncertainty, providing the setting in which verification behaviour and reliance can be examined more directly.

## 6. Expected Next Steps and Final Contribution to Knowledge

The remaining work focuses on resolving the behavioural consequences of the miscalibration pattern established in earlier stages. The attitudinal baseline is empirically supported, and interaction-level mechanisms have demonstrated technical feasibility. What remains uncertain is whether these mechanisms translate into observable differences in verification behaviour when entrepreneurs engage with AI systems in decision settings. The agentic tone calibration study examines whether tone calibration preserves perceived personal agency in agentic interaction. The GraphRAG rationale study evaluates whether graph-conditioned rationales increase evidence inspection and verification relative to provenance-only explanations. The entrepreneurial decision study tests whether the moderation pattern predicted in H8 emerges in an entrepreneurial decision scenario, and whether explainability reduces it. Together, these studies determine whether explanation and conversational design features improve calibrated reliance and preserve entrepreneurial agency, rather than merely making systems feel more trustworthy.

The contribution of this work is a mechanism-level account of attitudinal-behavioural miscalibration in AI-mediated decision support, specified through the conditions under which perceived trust diverges from verification behaviour. If supported, the findings have three implications. Perceptual trust measures alone are insufficient for evaluating explanation artefacts because they do not capture whether users engage in verification behaviour. An explanation that increases perceived understanding without changing verification behaviour does not meet the evaluative standard adopted here. Persona adaptation likewise cannot be treated as a neutral interface feature because, where it weakens verification behaviour, it reproduces the same trust inflation observed in the attitudinal baseline at the behavioural level. More broadly, AI interaction design functions as a situational determinant of entrepreneurial agency because it shapes whether actors evaluate, revise, and enact decisions through their own judgment or ratify a system-generated path. This implies that entrepreneurial decision-making research must explicitly account for AI interaction design as a condition that shapes agency, and that AI decision-support tools should be evaluated based on their ability to support informed decision behaviour rather than on

adoption alone. This is particularly relevant given the emerging co-agency literature, which documents capacity augmentation through AI but has not yet addressed whether such augmented capacity is accompanied by authorship over the resulting decisions [24].

## Acknowledgments

This research is part of the project LESSEN with project number NWA.1389.20.183 of the research program NWA\_ORC 2020/21, which is (partly) funded by the Dutch Research Council (NWO).

## Declaration on Generative AI

During the preparation of this work, the author used ChatGPT 5.4 (OpenAI) to assist with grammar checking and structural editing of draft text. After using this tool, the author reviewed and edited all content as needed and takes full responsibility for the publication's content.

## References

- [1] M. Fernández-Fernández, Á. Hernández-Tamurejo, P. González-Padilla, The ai trust factor: an mca analysis of automation and decision-making in entrepreneurship, *International Entrepreneurship and Management Journal* 22 (2026) 34. doi:10.1007/s11365-026-01170-4.
- [2] N. Scharowski, S. A. Perrig, N. von Felten, F. Brühlmann, Trust and reliance in XAI — distinguishing between attitudinal and behavioral measures, *arXiv preprint arXiv:2203.12318* (2022). doi:10.48550/arXiv.2203.12318.
- [3] A. Klingbeil, C. Grützner, P. Schreck, Trust and reliance on ai — an experimental study on the extent and costs of overreliance on ai, *Computers in Human Behavior* 160 (2024) 108352. doi:10.1016/j.chb.2024.108352.
- [4] J. Rady, D. Townsend, R. Hunt, From algorithmic hallucinations to alien minds: Addressing the ideator's dilemma through entrepreneurial work, *Journal of Business Venturing* 41 (2026) 106550. doi:10.1016/j.jbusvent.2025.106550.
- [5] F. de Véricourt, H. Gurkan, Is your machine better than you? you may never know, *Management Science* 72 (2026) 558–574. doi:10.1287/mnsc.2023.4791.
- [6] D. A. Shepherd, A. Majchrzak, Machines augmenting entrepreneurs: Opportunities (and threats) at the nexus of artificial intelligence and entrepreneurship, *Journal of Business Venturing* 37 (2022) 106227. doi:10.1016/j.jbusvent.2022.106227.
- [7] Y. Wang, K. Sauka, F. B. I. Situmeang, Anthropomorphism and transparency interplay on consumer behaviour in generative ai-driven marketing communication, *Journal of Consumer Marketing* 42 (2025) 512–536. doi:10.1108/JCM-04-2024-6806.
- [8] K. Sauka, Y. Saou, F. B. Situmeang, M. Kackovic, A nexus of explainability and anthropomorphism in AI chatbots, in: *World Conference on Explainable Artificial Intelligence*, volume 2576, Springer, 2025, pp. 115–138. doi:10.1007/978-3-032-08317-3\_6.
- [9] M. Emirbayer, A. Mische, What is agency?, *American Journal of Sociology* 103 (1998) 962–1023. doi:10.1086/231294.
- [10] J. Battilana, Agency and institutions: The enabling role of individuals' social position, *Organization* 13 (2006) 653–676. doi:10.1177/1350508406067008.
- [11] D. A. Shepherd, T. A. Williams, H. Patzelt, Thinking about entrepreneurial decision making: Review and research agenda, *Journal of Management* 41 (2015) 11–46. doi:10.1177/0149206314541153.
- [12] J. Schoeffer, M. De-Arteaga, N. Kühl, Explanations, fairness, and appropriate reliance in human-ai decision-making, in: *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, CHI '24, Association for Computing Machinery, New York, NY, USA, 2024. doi:10.1145/3613904.3642621.

- [13] H. Löfström, On the definition of appropriate trust and the tools that come with it, in: 2023 Congress in Computer Science, Computer Engineering, & Applied Computing (CSCE), 2023, pp. 1555–1562. doi:10.1109/CSCE60160.2023.00256.
- [14] M. Schemmer, N. Kuehl, C. Benz, A. Bartos, G. Satzger, Appropriate reliance on ai advice: Conceptualization and the effect of explanations, in: Proceedings of the 28th International Conference on Intelligent User Interfaces, IUI '23, Association for Computing Machinery, New York, NY, USA, 2023, p. 410–422. doi:10.1145/3581641.3584066.
- [15] Z. Buçinca, M. B. Malaya, K. Z. Gajos, To trust or to think: Cognitive forcing functions can reduce overreliance on ai in ai-assisted decision-making, *Proc. ACM Hum.-Comput. Interact.* 5 (2021). doi:10.1145/3449287.
- [16] F. Poursabzi-Sangdeh, D. G. Goldstein, J. M. Hofman, J. W. Wortman Vaughan, H. Wallach, Manipulating and measuring model interpretability, in: Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems, CHI '21, Association for Computing Machinery, New York, NY, USA, 2021. doi:10.1145/3411764.3445315.
- [17] R. Guidotti, A. Monreale, S. Ruggieri, F. Turini, F. Giannotti, D. Pedreschi, A survey of methods for explaining black box models, *ACM Comput. Surv.* 51 (2018). doi:10.1145/3236009.
- [18] A. Barredo Arrieta, N. Díaz-Rodríguez, J. Del Ser, A. Bennetot, S. Tabik, A. Barbado, S. Garcia, S. Gil-Lopez, D. Molina, R. Benjamins, R. Chatila, F. Herrera, Explainable artificial intelligence (xai): Concepts, taxonomies, opportunities and challenges toward responsible ai, *Information Fusion* 58 (2020) 82–115. doi:10.1016/j.inffus.2019.12.012.
- [19] E. Rajabi, K. Etminani, Knowledge-graph-based explainable ai: A systematic review, *Journal of Information Science* 50 (2024) 1019–1029. doi:10.1177/01655515221112844.
- [20] L. Longo, M. Brcic, F. Cabitza, J. Choi, R. Confalonieri, J. D. Ser, R. Guidotti, Y. Hayashi, F. Herrera, A. Holzinger, R. Jiang, H. Khosravi, F. Lecue, G. Malgieri, A. Páez, W. Samek, J. Schneider, T. Speith, S. Stumpf, Explainable artificial intelligence (xai) 2.0: A manifesto of open challenges and interdisciplinary research directions, *Information Fusion* 106 (2024) 102301. doi:10.1016/j.inffus.2024.102301.
- [21] L. Hondrich, H. Ruschmeier, Addressing automation bias through verifiability, in: EWAf 2023, European Workshop on Algorithmic Fairness, Proceedings of the 2nd European Workshop on Algorithmic Fairness, 2023.
- [22] D. Gursoy, O. H. Chi, L. Lu, R. Nunkoo, Consumers acceptance of artificially intelligent (ai) device use in service delivery, *International Journal of Information Management* 49 (2019) 157–169. doi:10.1016/j.ijinfomgt.2019.03.008.
- [23] S. D. Sarasvathy, Causation and effectuation: Toward a theoretical shift from economic inevitability to entrepreneurial contingency, *Academy of Management Review* 26 (2001) 243–263. doi:10.5465/amr.2001.4378020.
- [24] M. A. Al-Bashrawi, M. A. Al-Sharafi, I. A. Elgendy, M. Y. Helal, M. K. Anbalagan, I. Chae, Y. K. Dwivedi, Agentic ai systems and the future of entrepreneurship: a perspective on co-agency, innovation, and ecosystem transformation, *International Entrepreneurship and Management Journal* 22 (2026) 27. doi:10.1007/s11365-026-01164-2.
- [25] S. Murtinu, A. D. Massis, Artificial intelligence machines as relational nonhuman actors in entrepreneurial teams, *Journal of Small Business Management* 63 (2025) 2890–2916. doi:10.1080/00472778.2025.2461031.
- [26] J. J. Cai, X. Gu, L. Sheng, M. Xia, L. Zhao, W. Zhu, Ai as" co-founder": Genai for entrepreneurship, *arXiv preprint arXiv:2512.06506* (2025). doi:10.48550/arXiv.2512.06506.
- [27] D. M. Townsend, R. A. Hunt, Entrepreneurial action, creativity, & judgment in the age of artificial intelligence, *Journal of Business Venturing Insights* 11 (2019) e00126. doi:10.1016/j.jbvi.2019.e00126.
- [28] K. Manoorkar, K. Sauka, A. Palmigiano, N. Wijnberg, Interactive intent disambiguation using Dempster-Shafer theory, in: Proceedings of the ProActLLM Workshop at CIKM2025, 2025. Forthcoming.