

Artificial Trust for Decision-Making in Human-AI Teamwork Steps and Challenges

Centeio Jorge, Carolina; Jonker, Catholijn M.; Tielman, Myrthe L.

Publication date
2023

Document Version
Final published version

Published in
Workshops at the 2nd International Conference on Hybrid Human-Artificial Intelligence, HHAI-WS 2023

Citation (APA)

Centeio Jorge, C., Jonker, C. M., & Tielman, M. L. (2023). Artificial Trust for Decision-Making in Human-AI Teamwork: Steps and Challenges. In P. K. Murukannaiah , & T. Hirzle (Eds.), *Workshops at the 2nd International Conference on Hybrid Human-Artificial Intelligence, HHAI-WS 2023* (Vol. 3456, pp. 150-156). (CEUR Workshop Proceedings).

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Artificial Trust for Decision-Making in Human-AI Teamwork: Steps and Challenges

Carolina Centeio Jorge¹, Catholijn M. Jonker^{1,2} and Myrthe L. Tielman¹

¹*Interactive Intelligence, Delft University of Technology, Delft, Netherlands*

²*LIACS, Leiden University, Leiden, Netherlands*

Abstract

Human-AI teams count on both humans and artificial agents to work together collaboratively. In human-human teams, we use trust to make decisions. Similarly, our work explores how an AI can use trust (in human teammates) to make decisions while ensuring the team's goal and mitigating risks for the humans involved. We present the several steps and challenges towards the development of an artificial-trust-based decision-making model.

Keywords

artificial trust, human trustworthiness, human-AI teamwork

1. Introduction

Our work focuses on how an artificially intelligent agent (also referred to as AI) can understand its human teammates in a wide range of teams, from search and rescue to healthcare, cooking, etc. In particular, we explore how an AI can use artificial trust to predict and understand whether a human will do a certain task and, if so, how well. With artificial trust, the agent might be able to make informed decisions which will improve team efficiency and reduce risks. The main body of research on trust in human-AI teams is growing, see e.g., [1, 2, 3, 4, 5, 6], but does not address artificial trust in humans often. For research that touches upon artificial trust in humans, we refer to [7, 8, 9, 10, 11, 12, 13].

Artificial trust (a term described to the process of trust where the trustor is an AI [12]), similarly to natural trust (when the trustor is human), can be seen as a construct of aspects that are relevant when a human teammate is asked to collaborate on a certain task, such as ability, benevolence, integrity, capacity, preference, etc, see e.g. [14, 15]. However, it is still an open question how relevant these aspects are in human-AI teams. Most of these constructs come from human-human studies and need to be tested in scenarios where humans and AIs

HHAI-WS 2023: Workshops at the Second International Conference on Hybrid Human-Artificial Intelligence (HHAI), June 26–27, 2023, Munich, Germany

*Corresponding author.

✉ C.Jorge@tudelft.nl (C. Centeio Jorge); C.M.Jonker@tudelft.nl (C. M. Jonker); M.L.Tielman@tudelft.nl (M. L. Tielman)

🌐 <https://research.tudelft.nl/en/persons/carolina-centeio-jorge> (C. Centeio Jorge); <http://www.catholijnjonker.nl/> (C. M. Jonker); <https://ii.tudelft.nl/Myrthe/> (M. L. Tielman)

🆔 0000-0002-6937-5359 (C. Centeio Jorge); 0000-0003-4780-7461 (C. M. Jonker); 0000-0002-7826-5821 (M. L. Tielman)



© 2023 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

are teammates. On the other hand, the multi-agent systems (MAS) community has since long used beliefs of trust and trustworthiness for decision-making, see e.g., [16, 17, 18, 19]. Our work aims at computing trust beliefs for the agent, as in MAS literature, while having a human as trustee (the entity being trusted). This requires us to reach inspiration from social sciences and run user studies to explore and validate our possibilities. In this short paper, we present the steps and challenges towards our artificial-trust-based decision-making model.

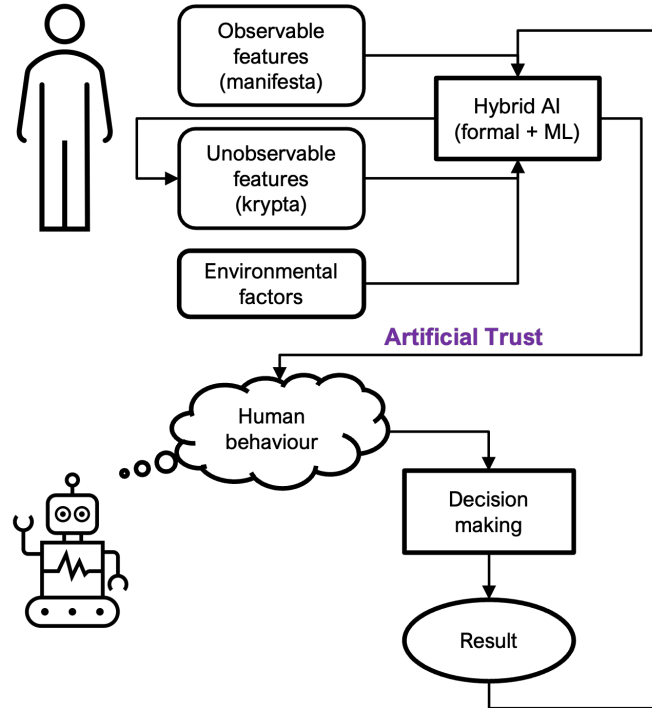


Figure 1: Overview of the model of artificial trust for decision-making. The AI system can observe the human teammate and, by estimating their features, model artificial trust in a hybrid way, i.e., using both knowledge and data-driven techniques. With beliefs of artificial trust, the AI system can estimate human behaviour and make a decision. Finally, it can update its model with the consequences of such decisions.

2. Towards beliefs of Artificial Trust in Human Teammates

The goal of this work is to enable an AI teammate to make decisions taking its human teammates' trustworthiness into account. For this, the AI teammate needs to be able to form beliefs of artificial trust regarding humans. In [20], the authors present an artificial agent's trust in another as a construct of *Competence belief*, *Willingness belief* and *Dependence belief*. Competence belief deals with believing the trustee has the necessary abilities to perform a task, whereas willingness translates into believing the trustee will do a task *given the context*, independently of their abilities. Finally, dependence belief lies on the trustor's side, and it is crucial for the decision-making process as it tells how much the trustor depends on the trustee for the execution of a

certain task. When we consider a team and joint goals, not only the trustor's (the one making the decision to trust) dependence belief is important but also the trustee's dependence on the trustor (e.g., the AI may choose to help a human because it believes the human is in a risky situation and depends on the AI for success). To the best of our knowledge, these beliefs are yet to be implemented and tested in human-AI interaction. It is challenging to do so as we do not know how to form such beliefs from humans, i.e., as said before, which aspects constitute human trustworthiness and how an artificial agent can perceive them. However, we aim at having such beliefs as a starting point for our artificial trust model, which will be used to make decisions.

Figure 1 presents the overview of the goal model. The human presents a *manifesta* (i.e., set of behavioural cues) which represent a *krypta* (i.e., set of characteristics of the human) [21, 22]. The AI can use the *manifesta* to model beliefs of artificial trust in a hybrid way, i.e., with both data-driven and knowledge-based techniques, from the *manifesta* along with *environmental factors* (which give context). Beliefs of artificial trust can be, as mentioned before, *competence*, *willingness*, and *dependence* beliefs. With some of these beliefs, such as competence and willingness, the AI should be able to predict some human behaviour and then, keeping the context in mind (and the *(inter)dependencies*), make a decision. Finally, given the outcome of this action, the agent should update its model of artificial trust.

In summary, the modelling steps towards a model of artificial trust for AI teammates' decision-making are:

1. *Investigating krypta and manifesta of human trustworthiness towards an artificial teammate.* This included exploring which unobservable characteristics (the *krypta*) constitute human trustworthiness in human-AI teams and how they can be observed (the *manifesta*) through a user study in an online 2D grid-world supermarket environment. The task consisted of helping two artificial agents by collecting the necessary products in the supermarket (inspired by the increase of online supermarket orders during the COVID-19 pandemic). The agents would ask for different products and the human could choose which agent to help. After choosing to help one of the agents, the subject could either complete the task, lie about it, or give up on that task, and then go to the next one. We included metrics and conditions based on Mayer's ABI model [14], which proposes that someone's trustworthiness depends on their ability, benevolence, and integrity. The results of the experiment are currently in the publishing process, but it is possible to find a preliminary report in [23].
2. *Updating artificial trust based on interaction given a context.* After studying which aspects may play a role in a human's trustworthiness towards an artificial teammate, and grasping how it may be possible to perceive them, it is important to explore how this trust can be updated. Trust is dynamic[24], and an artificial teammate should be able to constantly update its trust values throughout the interaction. Here, we can integrate existing models such as [11].
3. *Using artificial trust to make decisions.* The goal of this step is to propose a decision-making model for an AI agent which takes into account values of artificial trust for each sub-task of a joint goal and, at the same time, the interdependencies. We use principles of Coactive Design [25], including Interdependence Analysis.

4. *Evaluating artificial trust and decision-making models.* It is, in general, hard to evaluate artificial trust models, since we do not have ground truth. We work on defining metrics to compare different artificial trust models with baselines. The values of artificial trust can lead to decisions that may impact the environment and possibly the human teammate. It is then easier to compare the outcome of the decisions, rather than the trust values alone. The baselines include never-trusting models, always-trusting models, and random models.

3. Challenges

Although the steps towards having an artificial-trust-based decision-making model are defined, they present several challenges. Overall, designing and evaluating experiments presents methodological and theoretical issues, such as trying to explore complex real-world scenarios in controlled environments [26, 27]. Our research is no exception, and we have learned along the way about the main obstacles.

The first and main challenge is the lack of ground truth. Human trustworthiness has no available ground truth, making it hard to evaluate our artificial trust models. When we propose objective measures of trustworthiness, which may be manifestations of the *krypta*, we cannot prove that these are correct, or good, as there is nothing to compare them with. There is even the risk that, when we propose what human trustworthiness in a certain context may be, define the measures and then design the study, and all of this without other models to compare with or ground truth, we fall into a self-filling prophecy. We have tried to compare our measures with humans' perception of their own trustworthiness, but this is far from the ground truth, as their perception can be far from reality. This challenge affects the first step. This being said, we develop metrics and baselines (mentioned in step 4) that help us evaluate our model by comparison.

Another challenge every time we start designing an experiment within this topic is the task design. The task is the platform in which we want to explore several aspects, but finding one task that accommodates all aspects is hard. Furthermore, to study human-AI teamwork, we need to ensure the human participant understands the collaborative aspect of the environment. Otherwise, participants might focus on solving the task rather than caring about the interaction (which, we believe, would not necessarily be the same in a real-world situation). For example, one may present different narratives to different participants, such as the context or background story of the AI system or their (human and AI) relationship, without integrating them into the task. This can lead to the participants focusing mainly on completing the task and ignoring the rest of the information they were given. For this reason, we invested in visual representations of the teamwork and of the levels of interdependencies, with the aim that the participants understand that the AI is collaborative and that they can team up.

4. Conclusion

This paper presents an overview of the work that is being developed towards an artificial-trust-based decision-making model. The goal of such a model is to enable an AI teammate to

make informed decisions within a team context, taking into account its human teammate's trustworthiness and the context. The model should be updated throughout the interactions. We explain how we try to bridge multi-agent systems with social sciences and present the main steps that we take to develop such a model. Finally, we conclude with some of the challenges we have encountered throughout our research, mainly related to the user studies, and reflect on tentative mitigation strategies. Although there is still a long gap to fill, we believe this model will be important for human-AI teams, improving their efficiency and mitigating possible risks.

Acknowledgments

This research was supported by the Delft AI Initiative and by EU Horizon 2020 research and innovation programme under GA Numbers 952215 (TAILOR), and 820437 (Humane AI Net), and supported by the National Science Foundation (NSF) under Grant Numbers 024.004.022 (Hybrid Intelligence), 024.004.031 (ESDiT), and 024.005.017 (ALGOSOC). The support is gratefully acknowledged. Any opinions, findings, conclusions, or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the supporting organisations.

References

- [1] A.-S. Ulfert, E. Georganta, A model of team trust in human-agent teams, in: Companion Publication of the 2020 International Conference on Multimodal Interaction, ICMI '20 Companion, Association for Computing Machinery, New York, NY, USA, 2020, p. 171–176. doi:10.1145/3395035.3425959.
- [2] M. Lewis, H. Li, K. Sycara, Deep learning, transparency, and trust in human robot teamwork, in: Trust in Human-Robot Interaction, Elsevier, 2020, pp. 321–352.
- [3] K. E. Schaefer, B. S. Perelman, G. M. Gremillion, A. R. Marathe, J. S. Metcalfe, A roadmap for developing team trust metrics for human-autonomy teams, in: Trust in Human-Robot Interaction, Academic Press, 2021, pp. 261–300. doi:<https://doi.org/10.1016/B978-0-12-819472-0.00012-5>.
- [4] E. J. D. Visser, Marieke, M. M. Peeters, Malte, F. Jung, S. Kohn, Tyler, H. Shaw, R. Pak, M. A. Neerincx, Towards a theory of longitudinal trust calibration in human-robot teams, International Journal of Social Robotics 12 (2020) 459–478. doi:10.1007/s12369-019-00596-x.
- [5] B. G. Schelble, C. Lancaster, W. Duan, R. Mallick, N. J. McNeese, J. Lopez, The effect of AI teammate ethicality on trust outcomes and individual performance in human-ai teams, in: T. X. Bui (Ed.), 56th Hawaii International Conference on System Sciences, HICSS 2023, Maui, Hawaii, USA, January 3-6, 2023, ScholarSpace, 2023, pp. 322–331.
- [6] N. J. McNeese, M. Demir, E. K. Chiou, N. J. Cooke, Trust and team performance in human-autonomy teaming, Int. J. Electron. Commer. 25 (2021) 51–72. doi:10.1080/10864415.2021.1846854.
- [7] A. R. Wagner, R. C. Arkin, Recognizing situations that demand trust, in: 2011 RO-MAN, IEEE, 2011, pp. 7–14.
- [8] A. R. Wagner, P. Robinette, A. Howard, Modeling the human-robot trust phenomenon: A

conceptual framework based on risk, *ACM Transactions on Interactive Intelligent Systems* 8 (2018). doi:10.1145/3152890.

- [9] S. Vinanzi, M. Patacchiola, A. Chella, A. Cangelosi, Would a robot trust you? developmental robotics model of trust and theory of mind, *Philosophical Transactions of the Royal Society B: Biological Sciences* 374 (2019). doi:10.1098/rstb.2018.0032.
- [10] V. Surendran, A. Wagner, Your robot is watching: Using surface cues to evaluate the trustworthiness of human actions, 2019 28th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN) (2019) 1–8.
- [11] A. Ali, H. Azevedo-Sa, D. M. Tilbury, L. P. Robert, Heterogeneous human–robot task allocation based on artificial trust, *Scientific Reports* 12 (2022) 1–15.
- [12] H. Azevedo-Sa, X. J. Yang, L. P. Robert, D. M. Tilbury, A unified bi-directional model for natural and artificial trust in human-robot collaboration, *IEEE Robotics Autom. Lett.* 6 (2021) 5913–5920. doi:10.1109/LRA.2021.3088082.
- [13] S. Vinanzi, A. Cangelosi, C. Goerick, The collaborative mind: intention reading and trust in human-robot interaction, *iScience* 24 (2021) 102130. doi:10.1016/j.isci.2021.102130.
- [14] R. C. Mayer, J. H. Davis, F. D. Schoorman, An integrative model of organizational trust, *Source: The Academy of Management Review* 20 (1995) 709–734.
- [15] M. Johnson, J. M. Bradshaw, The role of interdependence in trust, in: *Trust in Human-Robot Interaction*, Elsevier, 2021, pp. 379–403.
- [16] J. Sabater-Mir, L. Vercouter, Trust and reputation in multiagent systems, *Multiagent systems* (2013) 381.
- [17] J. Urbano, A. P. Rocha, E. Oliveira, Computing confidence values: Does trust dynamics matter?, in: *Portuguese Conference on Artificial Intelligence*, Springer, 2009, pp. 520–531.
- [18] C. Castelfranchi, R. Falcone, Trust is much more than subjective probability: Mental components and sources of trust, in: *Proceedings of the 33rd annual Hawaii international conference on system sciences*, IEEE, 2000, pp. 10–pp.
- [19] R. Falcone, G. Pezzulo, C. Castelfranchi, A fuzzy approach to a belief-based trust computation, in: *Trust, Reputation, and Security: Theories and Practice*, volume 2631, Springer Verlag, 2003, pp. 73–86. doi:10.1007/3-540-36609-1_7.
- [20] R. Falcone, C. Castelfranchi, Trust dynamics: How trust is influenced by direct experiences and by trust itself, in: *AAMAS*, IEEE Computer Society, 2004, pp. 740–747. doi:10.1109/AAMAS.2004.10084.
- [21] R. Falcone, M. Piunti, M. Venanzi, C. Castelfranchi, From manifesta to krypta: The relevance of categories for trusting others, *ACM Transactions on Intelligent Systems and Technology* 4 (2013). doi:10.1145/2438653.2438662.
- [22] M. Bacharach, D. Gambetta, Trust as type detection, in: *Trust and deception in virtual societies*, Springer, 2001, pp. 1–26.
- [23] C. Centeio Jorge, M. L. Tielman, C. M. Jonker, Assessing artificial trust in human-agent teams: a conceptual model, in: C. Martinho, J. Dias, J. Campos, D. Heylen (Eds.), *IVA '22: ACM International Conference on Intelligent Virtual Agents*, Faro, Portugal, September 6 - 9, 2022, ACM, 2022, pp. 24:1–24:3. doi:10.1145/3514197.3549696.
- [24] L. Huang, N. J. Cooke, R. S. Gutzwiller, S. Berman, E. K. Chiou, M. Demir, W. Zhang, Distributed dynamic team trust in human, artificial intelligence, and robot teaming, in: *Trust in human-robot interaction*, Elsevier, 2021, pp. 301–319.

- [25] M. Johnson, J. M. Bradshaw, P. J. Feltovich, C. M. Jonker, M. B. Van Riemsdijk, M. Sierhuis, Coactive design: Designing support for interdependence in joint activity, *Journal of Human-Robot Interaction* 3 (2014) 43–69.
- [26] A. L. Brown, Design experiments: Theoretical and methodological challenges in creating complex interventions in classroom settings, *The journal of the learning sciences* 2 (1992) 141–178.
- [27] A. Collins, D. Joseph, K. Bielaczyc, Design research: Theoretical and methodological issues, *Journal of the Learning Sciences* 13 (2004) 15 – 42. doi:10.1207/s15327809jls1301_2, cited by: 1090.