

The Role of Trust in Decision-Making for Human Robot Collaboration

Min Chen^{*1}, Stefanos Nikolaidis^{†1}, Harold Soh^{*}, David Hsu^{*} and Siddhartha Srinivasa[†]

^{*}Department of Computer Science, National University of Singapore

[†]The Robotics Institute, Carnegie Mellon University

chenmin@comp.nus.edu.sg, snikolai@cmu.edu, harold@comp.nus.edu.sg,
dyhsu@comp.nus.edu.sg, siddh@cmu.edu

I. INTRODUCTION

Human trust is a key component for effective integration of autonomy. Research has shown that inappropriate levels of trust in the system can lead to over-reliance, or under-reliance with detrimental effects on performance [Lee and See, 2004].

Consider, for example, the table-clearing example of Fig. 1 where a human and a robot collaborate to clear a table of its items. When we asked users to perform this task with a robot that was fully capable of clearing objects, we observed many inexperienced participants prevented the robot from moving the glass cup. They justified their interventions, saying that they did not trust the robot and that letting the robot move the glass cup was too risky. Clearly, human trust in the robot directly affected the perception of risk [Siegrist, 2000] and consequently, the interaction.

We propose a computational model that *integrates human trust into robot decision making*. Because human trust is not directly observable, we model it as a latent variable in a partially observable Markov decision process (POMDP). This trust-POMDP relies on two sub-models: (i) a trust dynamics model, which captures the evolution of human trust in the robot, and (ii) a human decision-making model, which captures the probability of different human actions depending on trust. The POMDP formulation can accommodate a variety of trust models. We propose a data-driven approach, where we learn these two sub-models from data.

Although significant work has been done on real-time human trust elicitation and modeling [Lee and Moray, 1992; Floyd *et al.*, 2015; Xu and Dudek, 2015; Wang *et al.*, 2016], this work closes the loop between modeling trust and choosing robot actions to maximize team performance. Our model enables the robot to both infer and influence the collaborating human's level of trust. Returning to our table clearing example, our trust-POMDP robot first removes the three sealed water-bottles to develop trust, and only attempts to remove the glass cup at the end (Fig. 2). In contrast, a baseline robot that fails to consider trust removes the highest reward item (the glass



Fig. 1: Top: A robot and a human collaborate to clear a table. Bottom left: A robot attempts to grasp the glass cup first, causing the human teammate who does not trust the robot to intervene. Bottom right: A robot has a model of human trust and reasons with it for decision making. It increases human trust by picking up the three bottles first and then goes for the glass cup at the end, in order to minimize human intervention and save human effort.

cup) first, resulting in unnecessary interventions by human teammates with low initial trust.

We conducted an Amazon Mechanical Turk (AMT) study to compare our trust-POMDP against a baseline model that did not take human trust into account during decision-making. Experimental results substantiates our hypothesis that reasoning with human-trust improved team performance: the trust-POMDP was able to increase human trust when it was too low and significantly reduced the number of human interventions.

¹Authors contributed equally.

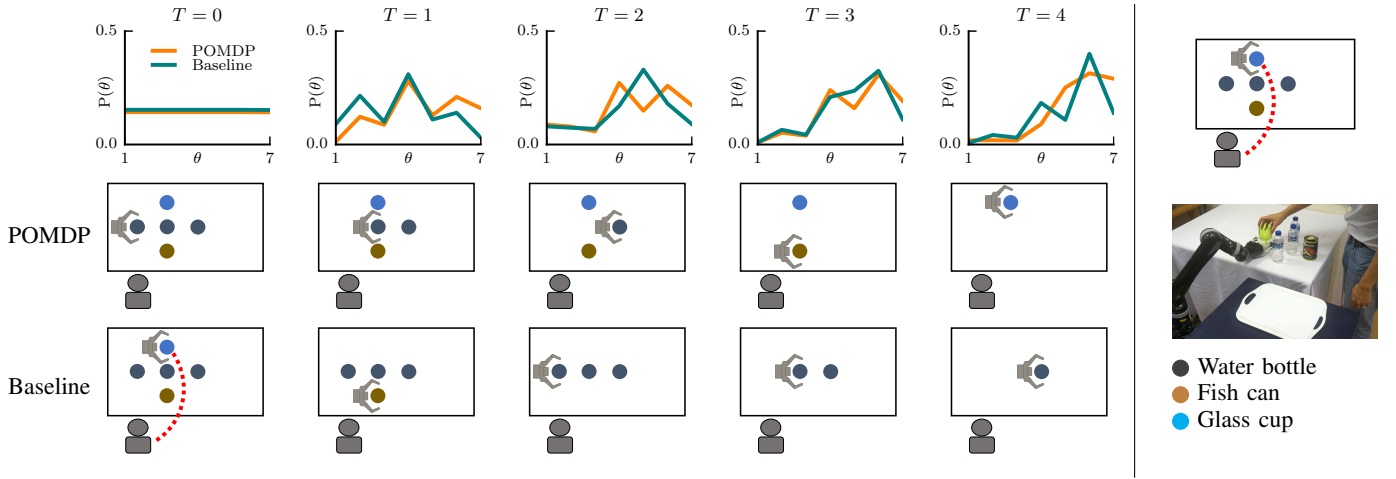


Fig. 2: Sample runs of the trust-POMDP (middle-row) and baseline (bottom-row) policies for the collaborative table-clearing task. The top row shows the probabilistic estimate of human trust θ over time (ignored by the baseline policy). Each pictogram represents a distinct snapshot of the task. The trust-POMDP policy starts with the bottles to build trust and chooses the glass cup only when the estimated trust is high enough. The baseline policy starts with the glass cup, causing the human to intervene.

II. PROBLEM FORMULATION

A human-robot team can be treated as a multi-agent system, with world state $x \in X$, robot action $a^R \in A^R$, and human action $a^H \in A^H$. The system evolves according to a stochastic state transition function $\mathcal{T}: X \times A^R \times A^H \rightarrow \Pi(X)$. At each time step, the human-robot team receives a real-valued reward $\mathcal{R}(x, a^R, a^H)$. We denote $h_t = \{x_0, a_0^R, a_0^H, x_1, \dots, x_{t-1}, a_{t-1}^R, a_{t-1}^H\} \in H_t$ as the history of interaction between robot and human until time step t . We assume that the human follows a stochastic policy, $\pi^H: X \times A^R \times H_t \rightarrow \Pi(A^H)$, that is unknown to the robot.

Our objective is to compute an optimal policy π^{R*} for the robot that maximizes the expected total discounted reward:

$$\pi^{R*} = \arg \max_{\pi^R} \mathbb{E}_{\pi^H} \left[\sum_{t=0}^{\infty} \gamma^t \mathcal{R}(x_t, a_t^R, a_t^H) | \pi^R, \pi^H \right] \quad (1)$$

where x_t, a_t^R, a_t^H denote the state, action taken by the robot and action taken by the human at time step t . The expectation is taken over the human behavioral policies π^H and the sequence of uncertain state transitions over time. For the robot to solve the optimization problem from Eq. 1, the robot needs access to the human behavioral policies π^H . In general, the human behavior may depend on the entire history of interactions h_t , which can grow arbitrarily large.

Our key insight is that in a number of human-robot collaboration scenarios, *trust is a compact approximation of the history of interactions h_t* . Therefore, we hypothesize that we can use trust as a predictor of future human actions. This allows us to condition human behavior on the inferred trust level, and in turn find the optimal policy that maximizes team performance.

However, trust cannot be directly observed by the robot and therefore, has to be inferred. Furthermore, trust in the robot is likely to change depending on the robot's performance.

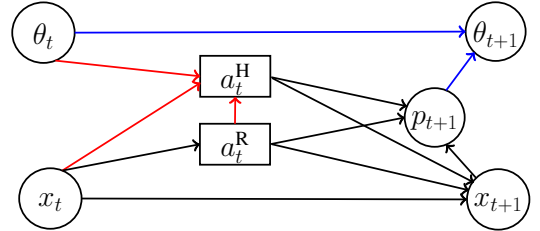


Fig. 3: State transition dynamics of trust-POMDP. The human action a_t^H at time t is governed by a stochastic policy $P(a_t^H | x_t, \theta_t, a_t^R)$ that depends on the world state x_t , the robot action a_t^R , and the current trust level θ_t (red arrows). The human and robot actions influence the next world state x_{t+1} . $p_{t+1} = f(a_t^H, a_t^R, x_{t+1})$ is the robot's performance at time t . Given the robot performance and the previous trust level, trust then updates stochastically via $P(\theta_{t+1} | \theta_t, p_{t+1})$ (blue arrows).

Both these issues can be addressed by the partially observable Markov decision process (POMDP), which provides a principled and general framework for sequential decision making under uncertainty.

Our trust-POMDP takes trust as a latent state variable θ , and includes a model of human behavior and trust dynamics into the state transition function (Fig. 3). The solution to a trust-POMDP is a policy that maximizes the team performance.

III. EVALUATION

We validated our trust-POMDP model via an online human subject experiment (201 participants) on the collaborative table-clearing task. Before the robot had reached each object, the video of the robot moving towards the object paused, and participants could choose to intervene and pick up the object themselves. Participants' intervention rate decreased by 31% and 54% when our trust-POMDP robot attempted to remove

the glass and can objects, compared to a baseline robot that did not account for trust, resulting in a statistically significant improvement in team performance.

We look forward to extend this evaluation with real-world human subject experiments, and to applying our trust-POMDP to enhance the usability of autonomous systems in a variety of collaborative settings.

REFERENCES

- Michael W Floyd, Michael Drinkwater, and David W Aha. Trust-Guided Behavior Adaptation Using Case-Based Reasoning. In *Proceedings of the Twenty-Fourth International Joint Conference on Artificial Intelligence*, pages 4261–4267, 2015.
- John Lee and Neville Moray. Trust, control strategies and allocation of function in human-machine systems. *Ergonomics*, 35(10):1243–1270, 1992.
- John D Lee and Katrina A See. Trust in automation: Designing for appropriate reliance. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 46(1):50–80, 2004.
- Michael Siegrist. The influence of trust and perceptions of risks and benefits on the acceptance of gene technology. *Risk analysis*, 20(2):195–204, 2000.
- Ning Wang, David V Pynadath, and Susan G Hill. Trust calibration within a human-robot team: Comparing automatically generated explanations. In *2016 11th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, pages 109–116. IEEE, 2016.
- Anqi Xu and Gregory Dudek. Optimo: Online probabilistic trust inference model for asymmetric human-robot collaborations. In *HRI*, pages 221–228. ACM, 2015.