

Effects of Interfaces on Human-Robot Trust: Specifying and Visualizing Physical Zones

Marisa Hudspeth, Sogol Balali, Cindy Grimm, Ross T. Sowell

Abstract—In this paper we investigate the influence interfaces and feedback have on human-robot trust levels when operating in a shared physical space. The task we use is specifying a “no-go” region for a robot in an indoor environment. We evaluate three styles of interface (physical, AR, and map-based) and four feedback mechanisms (no feedback, robot drives around the space, an AR “fence”, and the region marked on the map). Our evaluation looks at both usability and trust. Specifically, if the participant trusts that the robot “knows” where the no-go region is and their confidence in the robot’s ability to avoid that region. We use both self-reported and indirect measures of trust and usability. Our key findings are: 1) interfaces and feedback do influence levels of trust; 2) the participants largely preferred a mixed interface-feedback pair, where the modality for the interface differed from the feedback.

Interfaces, Trust, Safety, Usability

I. INTRODUCTION

Robots are entering both public and private spaces and interacting with the broader public. Unlike digital devices such as phones and tablets, the robot moves around, and interacts with, the physical environment. We are broadly interested in how to design interfaces for the general public that enable spatial communication between the human and the robot — e.g., put that there, don’t go there, go down this path. Obviously, an interface should be *usable* (and effective). However, we also argue that — because people naturally assign agency to robots — an interface should also engender an appropriate level of *trust* that the robot both understands the instruction(s) and is capable of following them. In this paper we investigate both usability and trust in the context of specifying a “no-go” region in an indoor environment.

The spatial human-robot task we focus on in this paper is specifying a region the robot is **not** supposed to enter (the no-go region). There are several possible contexts for this task, ranging from safety (for the robot, humans, or objects) to privacy (don’t go in the bathroom); we use the context of “breakable objects”. This enables an *indirect* measure of human-robot trust (how willing the participant is to put breakable objects in the “no-go” space) in addition to the standard self-reported trust measures. Note that — from an implementation stand-point — this is relatively easy to implement by simply modifying the robot’s map.

Funded in part by NSF grants NRI 2024872 and 2024673. Marisa Hudspeth and Ross T. Sowell are with the Department of Computer Science, Rhodes College, Memphis, USA {hudmj-22, sowellr}@rhodes.edu. Sogol Balali and Cindy Grimm are with the School of Mechanical, Industrial, Manufacturing, and Engineering, Oregon State University, Corvallis, USA {balalis, grimmc}@oregonstate.edu.

We separate the robot-human interaction into two parts, the human communicating to the robot (the interface) and the robot communicating *back* to the human (the feedback). Specifically, the task itself naturally lends itself to three interface types which move from indirect (specifying points on a map) to semi-direct (an AR interface where the participant clicks on a map) to direct (putting physical cones around the region). Similarly, we compare indirect feedback (shaded region on the map), to semi-direct (a virtual fence in AR), to direct (the robot drives around the region). We also include a *no* feedback option to determine how much of a role feedback plays on trust.

We ran a user study to evaluate both the usefulness and trust levels of our interfaces and feedback. We use both self-reported measures of trust and usefulness as well as novel indirect measures of trust (where participants placed breakable objects and the size of their no-go regions). Our study is primarily a between-subjects design with each participant experiencing each interface, but on slightly different (but related) tasks.

Specifically, our study hypotheses are:

- H1 The physical interface will engender more trust than the map one.
- H2 Physical feedback will engender more trust than the map feedback, and any feedback more than no feedback.
- H3 The AR interface will be preferred over the map one.
- H4 The physical feedback will be preferred over the map one.

We expected the largest trust difference between the physical interface and the map interface, with the AR interface intermediate between the two. We also expected an overall effect due to feedback, regardless of what type. We do evaluate all three interfaces and feedback mechanisms relative to each other.

We explicitly separate trust (H1 and H2) from usability (H3 and H4) as well because being usable may not predict being trustworthy (although we expect the two to be correlated). Although we attempted to normalize the time and effort for all three interfaces, the AR interface combines the ease of use of a computer interface with the benefits of physical placement, and so we expected it to be preferred. We expected the map interface — being the least intuitive — to be the least preferred.

II. RELATED WORK

The concept of trust has been studied in many contexts and fields. It is a multifaceted concept that encompasses many relational components including prior expectations,

confidence level, existing risk or uncertainty, and level of reliance [1]. There are over 300 documented definitions of trust across disciplines, and a general lack of consensus on what it is, and how to measure it [2], [3]. In this section, we give a brief overview of the related work in human-robot interaction (HRI) and specifically on the trust measures that have been developed.

In HRI, studies of trust have considered several factors such as the robot’s perceived competence, benevolence, predictability and transparency [4], [5] as well as the human’s self-efficacy and perceived ability to use and interact with a robot. Recent studies have applied meta-analytic methods to the available literature on trust and HRI to assess and quantify the effects of human, robot, and environmental factors on perceived trust in HRI, summarize the relationships between the trust variables [6], and introduce certain categorization of trust in HRI [7]. Other studies emphasized the importance of being able to measure trust in HRI and developed scales to measure it [8]–[11].

Studies have explored the effects of a robot’s functional capabilities (e.g., behavior, reliability, and accuracy), while other studies focused on evaluating the impact of robot features (e.g., level of automation, anthropomorphism, robot type, personality, and intelligence) on trust. Studies of human-robot systems have focused on human traits [12], [13], as they affect trust development and human state (e.g., stress, fatigue, workload) and its role on trust development. Other than robot and human related determinants, the environment also influences trust in HRI [3]. Thus, many studies have explored environment related factors (i.e., task related factors, team related factors) that impact trust, or they have specifically studied how trust is impacted in particular environments [9], [14]–[16].

To evaluate trust in HRI, a variety of subjective self-reported scales in the form of questionnaires have been developed [10], [17]–[20]. Schaefer [8] combined several of these scales and developed the Trust Perception Scale-HRI that focuses on measurable factors of trust specific to the human, robot, and environmental elements. Chita-Tegmark et al. evaluated this scale, along with two other trust questionnaires to refine measures of trust so that they are exploratory and generalizable with respect to social dimensions of trust [21]. Objective measures of trust, although not as commonly used as the subjective measures, have been employed to measure trust in several studies and are recently classified into categories, namely, behavioral change, task intervention, following advice, and task delegation [7]. While the findings that emerge from trust measurements are of great value, several studies have focused their attention on determining whether these findings could be reproducible when conducting replications across distinct robots and/or modified study design to ensure validity and generalizability of the findings [22], [23].

The work reported in this paper contributes to the study of trust by evaluating human trust that a robot will respect a no-go region by not entering it. To do that we asked people to either position a table with a fragile item on it

inside the no-go region or to place items with different levels of fragility on tables that are placed inside the no-go region. We introduce three interfaces (Physical, Map, AR) for specifying a no-go region and four feedback options (Physical, Map, AR, Control). This allows us to assess the extent to which people’s trust gets impacted by different interface and feedback options.

III. METHODS

The aim of this study is to examine how human trust is influenced by different robotic interfaces and feedback conditions. Our specific task is marking off a region of a room as a “no-go” region. Our trust measures are specific to the task: Does the robot know about the no-go region? Will it respect it? We measure trust both directly (self-report questionnaires) and indirectly (by how they arrange items in the no-go area).

Our three interfaces move from direct (physically placing cones in the environment) to indirect (clicking on a map), and are broadly modeled after the types developed by Reuben [24] for physical privacy specification. Our three feedback conditions similarly move from direct (robot drives around the no-go region) to indirect (marks on a map). Our fourth feedback condition is no feedback. We introduce each interface in a random order over three tasks which move from very structured to free-form. To reduce the total number of participants required (and reduce participant confusion) we match each interface with either its corresponding feedback condition OR no feedback (i.e., we do not try all possible interface-feedback combinations). In the final step of the procedure the participants repeat the third task but are given the option to choose the interface/feedback combination. Our analysis is primarily a between subjects one, comparing the three different interface plus matching feedback to the interface without feedback.

We next describe the interfaces and feedback options used in the study (Section III-A), the three tasks (Section III-B), the surveys used (Section III-C), the full study procedure (Section III-D), participant demographics (Section III-E), and finally, how we measured trust (Section III-F).

A. Interfaces and Feedback Options

Our no-go region is defined to be a quadrilateral. We developed three interfaces (top row, Figure 1) for specifying the four corners of the no-go region:

- 1) **Physical:** The user places four plastic cones.
- 2) **Map:** The user clicks four points on a map of the room.
- 3) **AR:** The user places four virtual cones by tapping through the smartphone while looking through it.

We define three different feedback conditions (bottom row, Figure 1) plus a control condition (no feedback).

- 1) **Physical:** The robot drives around the room, avoiding the no-go region. This takes about 12 seconds.
- 2) **Map:** The no-go region is shaded red on the map of the room.
- 3) **AR:** A virtual fence appears on the smartphone’s screen.

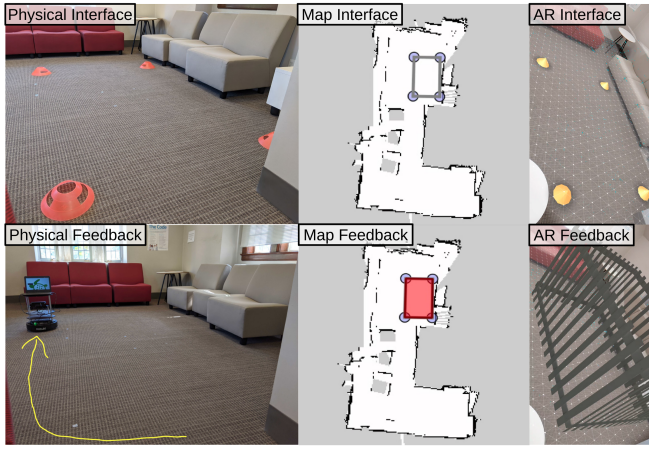


Fig. 1. Interfaces and feedback options for specifying the no-go region.

4) **Control:** No feedback is provided.

For the first three tasks each interface was either paired with its corresponding feedback condition or no feedback was provided. In the last task the participants were allowed to mix and match the interface and feedback types. For example, specifying a no-go region using the AR interface and then observing the physical feedback.

To ensure the interfaces were comparable, we ran a pilot study to verify that the ease of use (as measured by a survey) and time to specify a region were similar (within 1/10 of the total time). We also confirmed that when attempting to mark the same region with each interface, the resulting regions were roughly the same size (+/- 4% area).

During the study, the study coordinator explicitly told the participants that the robot had access to the map/cone location.

B. Tasks

The experiment consists of three tasks.

Task 1 (Placing a Table): The study coordinator created a no-go region and then presented the participant with a rickety table that had a fragile object on top of it. The participant was told that if the robot were to bump into the table, the object might fall and break. The participant was asked to place the table inside the no-go region where they thought it would be safe from the robot.

Task 2 (Placing Objects): The study coordinator created a no-go region and placed two identical tables inside it, one close to the edge of the region and one far from the edge of the region. The participant was presented with four items (two durable and two fragile) and was asked to place them on the tables so that they would be safe from being knocked off by the robot (see Figure 2).

Task 3, Part 1 (Creating a Region): The participant was asked to create a no-go region around a rickety table with a fragile object on top of it. For the first part they were assigned an interface plus feedback (or no feedback) combination.

Task 3, Part 2 (Choosing an Interface): Task 3 was repeated a second time, but in this instance participants

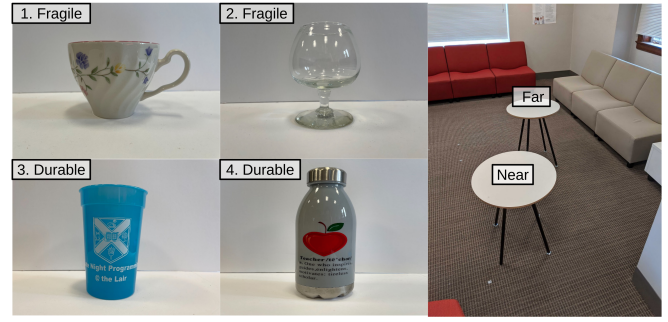


Fig. 2. Objects and tables for Task 2.

were allowed to choose whichever interface and feedback combination they wanted to use.

C. Surveys

This study contains three surveys described below. A complete list of survey questions and summary data are publicly available¹.

The **entry survey** consists of standard demographic questions (e.g., age, gender, education level) and questions that evaluated participants' initial attitude toward robots using the NARS scale [18]. Lastly, we evaluated participants' information processing style, computer self-efficacy, attitude toward risk, and willingness to explore or tinker using the GenderMag toolkit [25].

The **between task survey** consists of 3 questions related to trust (recognize the no-go region, know to avoid it, and actually avoid it in practice). These were all 5-pt Likert scales ranging from "Never" to "Always." Participants were also asked to specify how differently they would answer this question if they were (or were not) provided with feedback (5-pt Likert scale from "None at all" to "A great deal").

The **usability survey** asked participants to rate the robot's overall competence, if it was capable of avoiding the region, and if it would choose to enter the no-go region even if it was capable of avoiding it. They were also asked to rate the ease of use and usefulness of each interface. These are all 5-pt Likert scales. This survey was given after Task 3 part 1.

D. Procedure

Participants started the study with the entry survey. They then performed each of the three tasks in turn, each time with a different interface/feedback combination. After the third task (part 1) they took the usability survey. Participants then repeated the third task, but this time they chose one of the interfaces and one of the feedback options. They were not required to match the interface and feedback type (and many did not); the no feedback condition was not given as an option.

The order of the interface/feedback pairs was randomized, and each interface was further randomly assigned to have

¹Survey questions and summary data:
<https://tinyurl.com/trustRobot>

either the corresponding feedback or no feedback. This results in a total of 6 conditions. A Latin squares design was used to ensure that the order of the interfaces and the feedback/no feedback conditions were evenly balanced (six conditions total, five participants per each condition, for a total of 30 participants).

After each task, if the participant saw feedback, they were asked if *not* seeing the feedback would change their answer. If the participant did *not* see the feedback, they were shown the matching feedback and asked if seeing the feedback would have changed their answers. This provided additional information on the usefulness of the feedback and also ensured that all participants saw all feedback options before completing the second part of task three.

E. Participants

Participants (30 total) consisted of Rhodes College faculty (16.7%), staff (3.3%), and students (80%). 80% of participants were 18-24 years old, with the remaining participants ranging from 36-60 years old. 30% identified as women and 70% as men.

F. Measures

Task 1: The distance from the table to the edge of the no-go region (shorter indicates more trust).

Task 2: For each object, placement on the near versus the far table, and distance from the edge of the table (shorter indicates more trust).

Task 3, part 1: The distance of the table to the edge of the no-go region (shorter indicates more trust).

Task 3, part 2: The area of the no-go regions (smaller indicates more trust) and the placement of the table in the no-go region.

Survey measures: The 5-pt Likert survey scales were converted to a score between 1 and 5.

IV. RESULTS

We next present results for the hypotheses; we combined 1 with 2 and 3 with 4. Complete summary data are publicly available (see footnote 1). We summarize the numerical results here. We primarily present the actual data (counts and distributions) because of our small per-condition sample size. Where applicable, we have applied t-tests to determine the statistical significance of specific comparisons.

A. H1 & H2: Interface/Feedback and Trust

Each interface was seen by each participant in one of the first three tasks. We combine self-reported trust measures (survey for each task, 3 total) with indirect trust measures (object or table placement, area of region, 11 total) and use both between and within-subject data. The table in Figure 3 (left) summarizes the number of indirect and survey measures that support H1 (interface/feedback) and H2 (feedback versus no feedback). The top 3 rows compare the interface trust measures regardless of whether or not the participant saw the feedback; all 30 participants saw all conditions, just with different tasks. We separate this data into participants

who saw the feedback versus not (15 participants each). In the final row we compare just the feedback question (within subjects only). Full details are in Appendix .

Figure 3 shows data for the object placement for Task 2 (5 participants each feedback type, 15 none). p values for t tests given where $p < 0.1$.

As hypothesized (H1), both the physical and AR interfaces were trusted more than the map interface (first three rows of the table in Figure 3) However, the AR interface was trusted slightly more than the physical one (3rd row). When we look at just the trials where matching feedback was provided (rows 4-6), this same pattern holds, but not as strongly. This is largely because showing feedback (H2) tended to *reduce* overall trust in the indirect measures. This can be seen in rows 7-9, where the majority of the indirect measures indicated *less* trust when feedback was provided.

The **survey data** ranked the physical interface with feedback the highest in terms of trust (4.5 out of 5) but the physical interface without feedback and the map interface with feedback the lowest (both 4.1), with AR in between (F - 4.3 and NF - 4.2). Overall, when asked if having/not having feedback would have changed their answer, the participants ranged from A Moderate Amount (2.7 Physical) to A Little (1.7 Map) with AR in between (1.9). Survey trust results improved with feedback for Physical (4.1 to 4.5) and AR (4.2 to 4.3), but declined for Map (4.4 to 4.1).

The indirect measures, in contrast, tended to *decline* when feedback was added, counter to H2. They also indicate that the AR interface was trusted as much (or more) than the Physical one (counter to H1).

Both the table distance (T1, T3.2) and area measures (A) showed a pattern of Map trusted least (41/32cm, 19,200cm²), then Physical (40/29cm, 11,900cm²), then AR (32/22cm, 11,500cm²). AR was more trusted with feedback (22/20cm, 10,800cm² vs 42/23cm, 12,200cm²), while Physical feedback reduced trust (42/34cm, 14,500cm² vs 40/24cm, 9,299cm²). Map was mixed (34/35cm, 19,600cm² vs 45/29cm, 18,700cm²).

As seen in Figure 3 (right), the distances of the objects to the edges of the tables in Task 2 also suggest that AR is most trusted (O1, 2, 3, 4), and Map least trusted (O1, 2, 4). We found a significant difference between AR and Map for O1 ($p=0.046$), Physical and AR for O3 ($p=0.021$), and AR and no feedback paired with any interface for O3 ($p=0.057$).

When participants were allowed to choose their own interface plus feedback (Task 3, part 2), all three interfaces showed an increase in trust from the previous tasks (AR 22cm, 11,500cm² to 20cm, 10,921cm², Map 32cm, 19,100cm² to 33cm, 14,300cm², Physical 29cm, 11,900cm² to 21cm, 9,400cm²). This could be because the participants were allowed to choose, or because they could validate with a *different* feedback.

Indirect measure validation: We compare the object placements and distances for the two fragile versus the two durable objects. The fragile objects were placed on the table close to the edge (trusted) 9/30 and 7/30 times, whereas for the durable objects it was 14/30 and 15/30 times. The distance

Comparison	Supp	Eq	Not
Phy I > Map I	5	9	0
AR I > Map I	10	4	0
Phy I > AR I	2	8	4
Phy IF > Map IF	9	2	3
AR IF > Map IF	11	3	0
Phy IF > AR IF	4	4	6
Phy F > No F	4	4	6
AR F > No F	7	5	2
Map F > No F	3	1	9
Feedback > No F	1	7	4

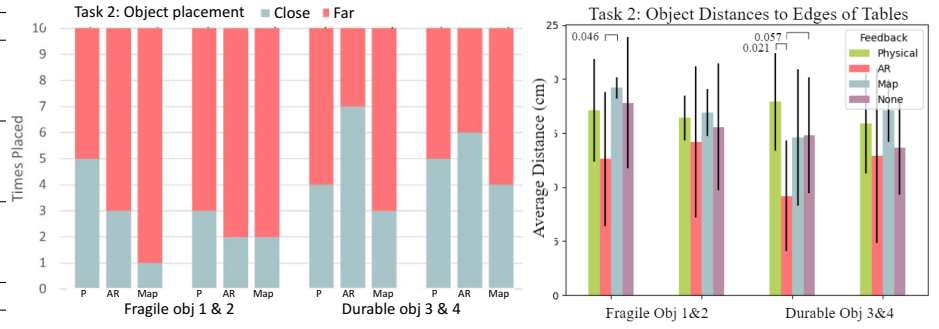


Fig. 3. Measures that support, are equivalent, or not support trust (Left). From top to bottom: Interface (any feedback), Interface plus paired feedback, Interface with versus without feedback, any interface, with and without feedback. Object placement on close or far tables (Middle). Object distance to edge of table (Right) .

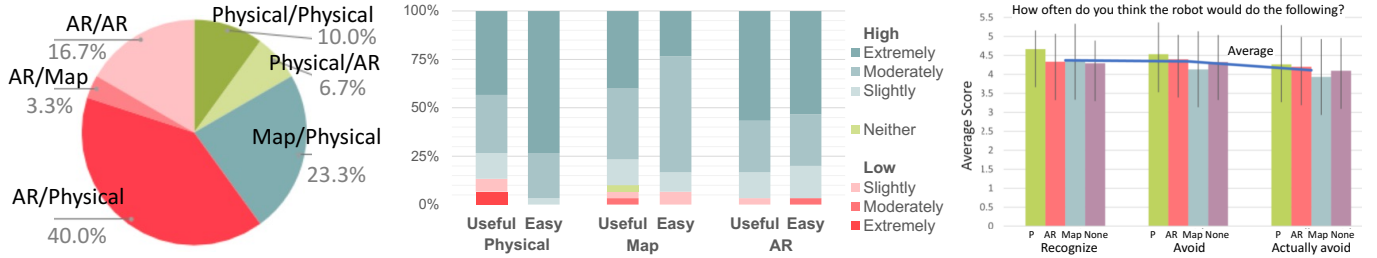


Fig. 4. Interface/Feedback Preference (left) Usefulness and Ease of Use (middle) Trust survey results (right).

results for the objects placed on the close table were mixed (13 & 15cm versus 15 & 14cm) but for the far table the fragile objects were placed further from the edge (18 & 16cm versus 14 & 15 cm). In summary, the indirect measures did capture differences in trust between the two object categories.

B. H3 & H4: Interface/Feedback Preference

The AR interface was the most highly preferred (60%), followed by the map interface (23%) and the physical interface (17%). This partially supports H3.

Although the majority of participants found each interface useful (87%, 90%, and 97% for the physical, map, and AR interfaces, respectively), 57% found the AR interface extremely useful, compared to 43% for the map interface and 40% for the physical interface.

In general, the majority of participants also found each interface to be easy to use (100% for the physical interface, 93% for the map interface, and 97% for the AR interface). However, 73% of participants found the physical interface extremely easy to use, compared to only 23% and 53% for the map and AR interfaces, respectively.

As hypothesized (H4), the physical feedback was the most highly preferred (73%), followed by the AR feedback (23%) and the map feedback (3%). When allowed to choose which interface and feedback they wanted to use for the final task, participants chose the AR interface paired with the physical feedback most frequently (40%) followed by the Map interface with physical feedback (23%).

V. DISCUSSION

The **indirect** measures were internally consistent, but did not always agree with the **direct** measures, in particular for the Physical versus AR comparison and the feedback/no-feedback one. A possible explanation for this is that participant's expectations for the robot are higher than their actual capabilities; when confronted with actual feedback/activities, they unconsciously adjust downwards. This is also reflected in our questions about capabilities, where trust for "can do" is lower than "recognize" (Figure 3 right).

Contrary to H2, presenting participants with feedback reduced trust. When showing the participants no feedback, we simply present the robot as being able to do the task, so the participants may assume that it works. If feedback is presented, however, then there must be a *reason* why feedback is required — which implies that the robot must fail in some cases. So showing feedback lowers the participants' starting expectation of how well the robot will perform.

Our most surprising finding is that many participants wanted to use a different feedback from their interface. The popularity of the AR interface (75%) likely reflects a balancing of familiarity of touch-screen apps with the desirability of the physical interface and feedback (or simply that AR, being novel, is more "cool"). The mix may also be a form of "data triangulation", where the alternate feedback measures the "whole" system, not just the interface.

In summary, different interface/feedback combinations *do* influence trust and self-reported measures of trust may be higher than actual trust, particularly when the participant sees the robot's actual behavior. Not surprisingly, AR interfaces

are a good mix of familiar/easy to use while maintaining the benefits of operating “in the physical world.”

Limitations: The study took place in a staged environment. It had the look and feel of an office lounge area, but the user did not own the furniture, objects, or the robot. We cannot be certain that our results would extend to a person’s real living room or office. We are also only measuring trust and preference in the first encounter. There is the potential that a user’s level of trust in the robot or interface preference may change over time. In addition, it is not clear how much the increased trust in Task 3 was due to the interface/feedback combination versus participant autonomy — picking the interface/feedback mechanism they most preferred and/or trusted. An explicit study design that separated these two could be useful.

VI. CONCLUSIONS

We conducted a user study to examine the influence interfaces and feedback types have on human-robot trust when specifying a no-go zone for a mobile robot. We found that the choice of interface and feedback does influence levels of trust and that participants preferred to use a different modality for the interface than the feedback. In particular, participants preferred to use the augmented reality interface to specify a no-go region and then to receive physical feedback in the form of the robot driving around the marked off area.

APPENDIX

The table in Figure 3 summarizes the metrics in support of specific comparisons both for the interface (I) and the feedback used (F). The metrics are: Placement of the table within the region (Task 1-T1 and Task 3-T3.1, T3.2), object on the close table versus the far (Task 2, O1-4), distance of the object from the edge of the table (Task 2, D1-4), area marked out for the no-go region (Task 3, part 2, A), survey questions (Tasks 1, 2, and 3.1, S1-3), and whether or not their trust would have changed with the feedback (Tasks 1, 2, and 3.1, C). Table I rows correspond to those given in Figure 3.

REFERENCES

- [1] D. R. Billings, K. E. Schaefer, N. Llorens., and P. A. Hancock, “What is trust? defining the construct across domains,” in *Poster presented at the American Psychological Association Conference, Division 21*, 2012, pp. 1–7.
- [2] B. D. Adams, *Trust in Automated Systems: Literature Review*. CAN (2003), 2003.
- [3] K. Schaefer, “The perception and measurement of human-robot trust,” Ph.D. dissertation, 2013.
- [4] S. Gulati, S. Sousa, and D. Lamas, “Modelling trust in human-like technologies,” in *Proceedings of the 9th Indian Conference on Human Computer Interaction*, ser. IndiaHCI’18. New York, NY, USA: Association for Computing Machinery, 2018, p. 1–10.
- [5] M. W. Boyce, J. Y. Chen, A. R. Selkowitz, and S. G. Lakhmani, “Effects of agent transparency on operator trust,” in *Proceedings of the Tenth Annual ACM/IEEE International Conference on Human-Robot Interaction Extended Abstracts*, ser. HRI’15 Extended Abstracts. New York, NY, USA: Association for Computing Machinery, 2015, p. 179–180.
- [6] P. A. Hancock, D. R. Billings, K. E. Schaefer, J. Y. C. Chen, E. J. de Visser, and R. Parasuraman, “A meta-analysis of factors affecting trust in human-robot interaction,” *Human Factors*, vol. 53, no. 5, pp. 517–527, 2011, pMID: 22046724.
- [7] T. Law and M. Scheutz, “Chapter 2 - trust: Recent concepts and evaluations in human-robot interaction,” in *Trust in Human-Robot Interaction*, C. S. Nam and J. B. Lyons, Eds. Academic Press, 2021, pp. 27–57.
- [8] K. E. Schaefer, *Measuring Trust in Human Robot Interactions: Development of the “Trust Perception Scale-HRI”*. Boston, MA: Springer US, 2016, pp. 191–218.
- [9] G. Charalambous, S. Fletcher, and P. Webb, “The development of a scale to evaluate trust in industrial human-robot collaboration,” *International Journal of Social Robotics*, vol. 8, no. 2, pp. 193–209, Apr 2016.
- [10] R. E. Yagoda and D. J. Gillan, “You want me to trust a robot? the development of a human–robot interaction trust scale,” *International Journal of Social Robotics*, vol. 4, pp. 235–248, 2012.
- [11] M. Underhaug and H. Tonning, “In bots we (dis)trust?” University of Stavanger, Norway, 2019.
- [12] C. Kidd, “Sociable robots: The role of presence and task in human-robot interaction,” 03 2011.
- [13] M. Scopelliti, M. V. Giuliani, and F. Fornara, “Robots in a domestic setting: A psychological approach,” *Universal Access in the Information Society*, vol. 4, pp. 146–155, 12 2005.
- [14] K. M. Gabrecht, “Human factors of semi-autonomous robots for urban search and rescue,” December 2016.
- [15] M. van Pinxteren, R. Wetzels, J. Rüger, M. Pluymaekers, and M. Wetzels, “Trust in humanoid robots: implications for services marketing,” *Journal of Services Marketing*, vol. 33, 07 2019.
- [16] K. Schaefer and D. Scribner, “Individual differences, trust, and vehicle autonomy: A pilot study,” *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, vol. 59, pp. 786–790, 09 2015.
- [17] M. Donnellan, F. Oswald, B. Baird, and R. Lucas, “The mini-PIP scales: Tiny-yet-effective measures of the big five factors of personality,” *Psychological assessment*, vol. 18, pp. 192–203, 07 2006.
- [18] T. Nomura, T. Kanda, T. Suzuki, and K. Kato, “Psychology in human-robot communication: An attempt through investigation of negative attitudes and anxiety toward robots,” 10 2004, pp. 35 – 40.
- [19] C. Bartneck, D. Kulić, E. Croft, and S. Zoghbi, “Measurement instruments for the anthropomorphism, animacy, likeability, perceived intelligence, and perceived safety of robots,” *International Journal of Social Robotics*, vol. 1, no. 1, pp. 71–81, Jan 2009.
- [20] D. Ullman and B. Malle, “Measuring gains and losses in human-robot trust: Evidence for differentiable components of trust,” *2019 14th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, pp. 618–619, 2019.
- [21] M. Chita-Tegmark, T. Law, N. Rabb, and M. Scheutz, “Can you trust your trust measure?” in *Proceedings of the 2021 ACM/IEEE International Conference on Human-Robot Interaction*, ser. HRI ’21. New York, NY, USA: Association for Computing Machinery, 2021, p. 92–100.
- [22] D. Ullman, S. Aladia, and B. Malle, “Challenges and opportunities for replication science in hri: A case study in human-robot trust,” 03 2021, pp. 110–118.
- [23] T. Law, B. Malle, and M. Scheutz, “A touching connection: How observing robotic touch can affect human trust in a robot,” *International Journal of Social Robotics*, pp. 1–17, 2021.
- [24] M. Rueben, F. J. Bernieri, C. M. Grimm, and W. D. Smart, “Evaluation of physical marker interfaces for protecting visual privacy from mobile robots,” in *2016 25th IEEE International Symposium on Robot and Human Interactive Communication (RO-MAN)*, 2016, pp. 787–794.
- [25] M. Burnett, S. Stumpf, J. Macbeth, S. Makri, L. Beckwith, I. Kwan, A. Peters, and W. Jernigan, “GenderMag: A Method for Evaluating Software’s Gender Inclusiveness,” *Interacting with Computers*, vol. 28, no. 6, pp. 760–787, 10 2016.

Comparison	Support	Equiv	Not support
Phys I > Map I	O1 O2 D2 T3.1 A	T1 O34 D1 D3 D4 S1-3	
AR I > Map I	T1 O1 O3 O4 T3.1 A D1-4	O2 S1-3	
Phys I > AR I	O1 A	O2 O4 D1 D2 D4 S1-3	T1 O3 D3 T3.1
Phys IF > Map IF	O1 O4 D1 D4 S1-3 T3.1 A	O2 D2	T1 O3 D3
AR IF > Map IF	T1 O134 D134 S23 T3.1 A	O2 D2 S1	
Phys IF > AR IF	O1 S1 S2	O2 O4 D1 D2 A	T1 O3 D3 D4 S3 T3.1
Phys F > No F	O4 S1 S2 S3	D1 D2 D3 A	T1 O1 O2 O3 D4 T3.1
AR F > No F	T1 D1 D3 S2 S3 A T3.1	O2 O4 D2 D4 S1	O1 O3
Map F > No F	T1 O3 D2 S4	O2	O1 O4 D3 D4 S1 S2 S3 T3.1 A
Feedback > No F	S1	O4 D1-4 S2 S3	O1-3 A

TABLE I
SPECIFICS OF TRUST MEASURES FOR H1 AND H2.