

Privacy and utility perceptions of social robots in healthcare

Sandhya Jayaraman^{a,*}, Elizabeth K. Phillips^b, Daisy Church^c, Laurel D. Riek^a

^a Department of Computer Science and Engineering, UC San Diego, 9500 Gilman Dr, La Jolla, 92093, California, USA

^b Department of Psychology, George Mason University, 4400 University Dr, Fairfax, 22030, Virginia, USA

^c Academy of Art University, 180 New Montgomery St, San Francisco, 94105, California, USA

ARTICLE INFO

Keywords:

Human-robot interaction
Privacy
Utility
Healthcare robots

ABSTRACT

Many Human-Robot Interaction (HRI) researchers are exploring the use of healthcare robots. Due to the sensitive nature of care, privacy concerns play a significant role in determining robot utility and adoption. While HRI research has explored some dimensions of privacy for robots in general, to our knowledge, no prior work has empirically studied how human-like robot design affects people's privacy and utility perceptions of robots across different healthcare contexts and tasks. We conducted a $3 \times 3 \times 3$ study ($n = 239$) to understand these relationships, varying robot Human Likeness (HL) (low, medium, and high) and scenario/task type (hospital waiting room/robot check-in support, hospital patient room/robot mobility support, home care/robot neuro-rehabilitation support) via a mixed between-within subjects design. To our knowledge, this is one of the first studies that operationalizes complex constructs of privacy, healthcare, and HL across multiple realistic healthcare contexts, with a high degree of cognitive fidelity. Our results suggest the tasks and contexts in which privacy is considered in healthcare contexts with robots is more impactful than other factors like robot HL appearance. In particular, some settings include more complex tradeoffs between privacy and utility for robots than others. For example, HRI researchers and practitioners who want to build healthcare robots intended for the home may encounter the greatest challenges for balancing privacy risks. Finally, for the community, we demonstrate that design fiction animations can be a useful way to facilitate cognitive fidelity for supporting studies in HRI and serving as a bridge between narrative methods and the use of real-world robots.

1. Introduction

Robots are entering people's daily lives, where they actively interact as social actors [Belpaeme et al. \(2018\)](#); [Leite et al. \(2013\)](#); [Broekens et al. \(2009\)](#). This increase in social robot adoption in human spaces, ranging from industrial environments to hospitals and to homes, has prompted vast interdisciplinary research on psychological [De Graaf and Allouch \(2013\)](#), ethical [Malle and Scheutz \(2020\)](#), design [Moharana et al. \(2019\)](#); [Taylor et al. \(2022\)](#), privacy [Rueben et al. \(2018, 2017\)](#) and technological considerations [Kubota et al. \(2020\)](#); [Alonso-Martín and Salichs \(2011\)](#) in developing such robots.

Healthcare, in particular, is witnessing a rapid growth in social robot adoption [Riek \(2017\)](#). Robots are increasingly serving in patient-facing roles, both in hospitals and homes, to support care delivery. In hospitals, especially during the pandemic, robots supported patient triage and check-in, delivered items, provided telemedicine, and companionship [Shen et al. \(2020\)](#). In homes, robots are supporting people with physical and cognitive rehabilitation, medication management, and physical task

assistance [Robinson et al. \(2014\)](#).

Something unique about the use of robots in healthcare is that they are providing support to people who may be in a vulnerable state. Like healthcare providers, healthcare robots may also have a duty to protect humans physically, psychologically, and socially. Many factors affect how people trust robots in healthcare, including concerns about their privacy, social influence from others, and familiarity with technology [Xu et al. \(2018\)](#); [Langer et al. \(2019\)](#); [Borenstein et al. \(2017\)](#). All of these factors ultimately impact robot adoption. Studies have shown that privacy concerns negatively impact people's trust in robots, and adversely affect their adoption [Alaiad and Zhou \(2014\)](#). Other studies have shown that perceived anthropomorphism of the robot positively correlates with people's trust in robots [Natarajan and Gombolay \(2020\)](#). This suggests HL could affect people's privacy perceptions of social robots.

Privacy-sensitive robotics is an emerging area of research that explores the unique privacy requirements of the embodied intelligent technologies that are robots. Prior research has characterized the dimensions of privacy associated with social robots [Lutz et al. \(2019\)](#),

* Corresponding author.

E-mail address: sjayaraman@ucsd.edu (S. Jayaraman).

<https://doi.org/10.1016/j.chbah.2023.100039>

Received 11 July 2023; Received in revised form 14 December 2023; Accepted 16 December 2023

Available online 25 December 2023

2949-8821/© 2024 The Authors. Published by Elsevier Inc. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

explored future research directions for privacy-sensitive robots [Rueben et al. \(2018\)](#), and exploited software and data vulnerabilities in existing robots [Denning et al. \(2009\)](#). However, many of these studies are qualitative with small sample sizes, and their results may not be generalizable to larger populations.

Anthropomorphism is the human tendency to attribute human-like characteristics to nonhuman objects. On the other hand, HL of robots is in part defined by the constituent human-like features that make up its overall design, where specific features and combinations of features predict how human-like the robot is perceived to be [Phillips et al. \(2018\)](#). Studies suggest that human-like design cues can have a positive effect on the acceptance of social robots, and in building long-term relationships with social robots [Fink \(2012\)](#).

A lot of prior privacy literature mentions the possibility of anthropomorphic robot design influencing people's perceptions of privacy [Lutz et al. \(2019\)](#); [Lutz and Tamó-Larrieux \(2020\)](#); [Rueben and Smart \(2016\)](#); [Lutz and Tamó \(2016\)](#); [Darling \(2015\)](#). However, to our knowledge, no prior research has empirically studied whether and how HL design affects people's privacy perceptions of robots across different healthcare contexts and tasks. Also, prior research confirms the existence of a privacy-utility tradeoff observed in human intent to use a social robot despite being aware of the privacy concerns implicitly associated with using it [Lutz and Tamó-Larrieux \(2020\)](#). We are interested to explore whether robot HL could influence privacy-utility trade offs.

In this paper, we explore three research questions related to human privacy perceptions of social healthcare robots.

RQ1: How is human-like robot design related to privacy perceptions?

RQ2: How is human-like robot design related to the privacy-utility tradeoff?

RQ3: How does context affect privacy perceptions of social healthcare robots?

To explore these questions, we conducted a large-scale quantitative study ($n = 239$) to understand how robot HL affects people's perceptions of privacy within the context of healthcare robots. We developed design fictions to serve as our stimuli, which were animated videos that depicted three social healthcare robots (Moro, Pepper, and Geminoid) providing assistance across three realistic healthcare contexts (hospital waiting room, hospital patient room, and home).

To our knowledge, this is one of the first studies that operationalizes complex constructs of privacy, utility, healthcare, and HL across multiple realistic healthcare contexts, with a high degree of cognitive fidelity. Most prior work in this space, including for privacy and social healthcare robots, employ static stimuli or vignettes, whereas we created a series of engaging design fiction [Noortman et al. \(2019\)](#); [Wong and Mulligan \(2016\)](#) animations, enabling participants to imagine future robots.

Our results revealed the tasks and contexts in which privacy is considered in healthcare contexts with robots is more impactful than other factors like robot human-like appearance. In particular, some settings include more complex tradeoffs between privacy and utility for robots than others. For example, HRI researchers and practitioners who want to build healthcare robots intended for the home may encounter the greatest challenges for balancing privacy risks.

2. Background

2.1. Privacy in HRI

Various definitions of privacy have been explored in HRI. [Rueben and Smart \(2016\)](#) discussed multidisciplinary definitions of privacy under broad themes of informational privacy, constitutional privacy, and access privacy. In their characterization of privacy themes for social robots, [Lutz et al. \(2019\)](#) also provided a number of existing definitions of privacy under four themes: informational privacy, social privacy,

psychological privacy, and physical privacy. Informational privacy refers to the privacy of personal information, social privacy refers to privacy among social actors, psychological privacy refers to the privacy of thoughts and values, and physical privacy refers to the privacy of physical boundaries.

While there is a large body of research on informational privacy, including legal and technical frameworks to preserve it, [Tavani \(2008\)](#); [Cohen \(2017\)](#); [Floridi \(2006\)](#), and application-specific considerations (e.g., drones [Luppici and So \(2016\)](#), the Internet of Things (IoT) [Lee \(2020\)](#)), informational privacy is insufficient for HRI. Social robots are embodied social entities, so exploring these other privacy concerns raised by [Lutz et al. \(2019\)](#) is equally important for HRI.

Generally HRI research explores privacy from a multifaceted lens. There are studies that explore the privacy implications of robots developed for a particular context. For example, Lutz et al. characterized privacy concerns of social robots in a scoping literature review with expert interviews [Lutz et al. \(2019\)](#). Other researchers have studied privacy with respect to telepresence robots [Rueben et al. \(2017\)](#); [Krupp et al. \(2017\)](#), healthcare robots [Lutz and Tamó \(2016\)](#), and household robots [Denning et al. \(2009\)](#). All such HRI studies acknowledge that defining privacy for robotics can be challenging, and address the importance of clearly stating the specific aspects of privacy that a study intends to address.

Privacy literature also notes the existence of discrepancies between user attitudes and user behaviors termed as the "privacy paradox" [Barth and De Jong \(2017\)](#). In other words, while users claim to be concerned to a certain level about their privacy, they do not follow the required measures to maintain this level of privacy they claim. HRI researchers have observed the privacy paradox in the context of social robot use, where perceived benefits of these robots outweigh the potential privacy concerns of using them [Lutz and Tamó-Larrieux \(2020\)](#). This inverse relationship between privacy concerns and perceived benefits has also been referred to as the privacy-utility tradeoff. Some HRI and Human-Computer Interaction (HCI) studies have explored privacy-utility tradeoffs through algorithmic approaches that enable data intensive applications to perform adequately and provide utility using minimal data [Butler et al. \(2015\)](#); [Jin et al. \(2022\)](#).

Empirical studies have used and developed different instruments to measure privacy. The Internet Users' Information Privacy Concerns (IUIPC) [Malhotra et al. \(2004\)](#) scale is one of the most widely used scales of information privacy. Other studies have adapted this scale to measure information privacy concerns in various different applications [Lutz and Tamó-Larrieux \(2020\)](#); [Dang et al. \(2021\)](#). Psychological scales of privacy measure privacy under the idea of the "right to be left alone" [Pedersen \(1999\)](#). Lutz and Tamó-Larrieux developed a measure of privacy for social robots which included informational privacy concerns, trusting beliefs, overall concerns, and developed a new subscale for measuring physical privacy [Lutz and Tamó-Larrieux \(2020\)](#).

2.2. HL in HRI

Robot appearance has been found to influence people's perceptions of a robot's intelligence [Haring et al. \(2016\)](#); [Sims et al. \(2005\)](#), credibility [Burgoon et al. \(2000\)](#), trust in the robot [Natarajan and Gombolay \(2020\)](#), and acceptance of the robot [Murphy et al. \(2019\)](#). Human-like appearance of robots can influence people's empathy towards the robot [Riek et al. \(2009\)](#), and willingness to work alongside the robot [Hancock et al. \(2011\)](#). Additionally, humans respond positively to human-like social cues, which can impact human-robot interaction and inform users' judgements of these robots [Hegel et al. \(2011\)](#); [Eyssel et al. \(2010\)](#).

Research has characterized features and dimensions of robots that contribute to people's perceptions of robot humanness [DiSalvo et al. \(2002\)](#); [von der Pütten and Krämer \(2012\)](#). Phillips et al. developed a measure of HL of robots [Phillips et al. \(2018\)](#) based on their constituent human-like physical features. Other works characterize

anthropomorphism as a combination of HL in appearance and HL in interaction and behavior [Carpinella et al. \(2017\)](#); [Spatola et al. \(2021\)](#).

3. Methodology

The study followed a 3 (HL: Low, Medium, High) $\times$ 3 (Scenarios/Robot Tasks: hospital waiting room/robot check-in support, hospital patient room/robot mobility support, home care/robot neuro-rehabilitation support) $\times$ 3 (repeated administration of measures) mixed between-within subjects design. The animated medical scenarios were treated as a within-subjects variable where participants viewed all three of the animated medical scenarios but with only one of the robots (i.e., between-subjects) depicted across each of the scenarios. We selected this design so participants did not need to frame switch between both different scenarios and different robots, as this can lead to participants getting cognitively overloaded.

3.1. Participants and power analysis

A power analysis was conducted using G*Power software [Faul et al. \(2007\)](#) for a mixed within-between subjects F-test to detect potentially small effect sizes with power of $\beta = 0.80$ and $\alpha = 0.05$. We also planned to sample 120% of this number to account for potential participant attrition in online studies. Thus, we aimed to recruit $N = 240$ participants for this study. Participants were recruited using Prolific. 239 participants (114 females, 118 males, 1 transgender, 4 non-binary, and 1 not reported), with ages ranging from 18 to 83 years, $M_{age} = 37.03$, $SD_{age} = 13.80$ completed the study. All participants passed “bot check” and audio-video check procedures before viewing our video stimuli.

3.2. Measures

We measured privacy perceptions using three subscales from the Lutz and Tamó-Larrieux privacy questionnaire for social robots [Lutz and Tamó-Larrieux \(2020\)](#) – trusting beliefs, overall privacy concerns, and physical privacy concerns. To measure informational privacy concern, we used the health information disclosure scale [Dang et al. \(2021\)](#), since it specifically addressed informational privacy from a health information perspective.

Additionally, we created a custom three-item utility scale to explore the privacy/utility tradeoff more explicitly. We first looked in the literature to explore existing perceptions of utility measures. We found that in many studies, researchers used existing measures of usability as a proxy for utility (e.g., [Klow et al. \(2017\)](#)), while others integrated perceived value into their measurement techniques (e.g., [Tran and Nguyen \(2021\)](#)). By definition (Merriam-Webster), utility differs from usability in that utility implies there is a specific need for the item/product/design independent of whether it works well, is liked, or is unnecessarily complex. Thus, existing usability measures are likely an insufficient representation of the construct of utility and the trade-off between privacy and utility. Thus, we derived three items based on the definitions given of utility that represent elements of the focal construct including being useful, beneficial, and fulfilling a need. To capture the privacy-utility tradeoff when administering the items after each healthcare scenario, we asked participants to imagine the ambiguous action of the robot and then respond to the item, e.g., “If the robot in this scenario did not clear its screen before the next check-in it could still [be useful/be beneficial/fulfill a need] for hospital check-in.” Participants provided their response using a 5-point Likert-type scale anchored from 1 (Strongly disagree) to 5 (Strongly agree). All our questions can be found in [Table 4](#).

We also asked participants to respond to one qualitative open-ended question at the end of the study: “Please use this space to leave us honest feedback concerning the study.”

3.3. Scenarios

To support ecological validity, we developed three scenarios depicting healthcare robots across three of the most common uses for healthcare robots (aside from surgical robots) [Kyrarini et al. \(2021\)](#); [Riek \(2017\)](#), each conveying dimensions of privacy risk. Each scenario intentionally ended in an ambiguous way, where it was not clear what the robot would do. This was to further explore the interactions between context, HL, and privacy concerns.

The first scenario was a hospital waiting room, and depicted a hospital check-in robot that asked triage questions. At the end of the scenario the robot had highly personal medical information still on the screen (e.g., the patient’s loss of bladder control) as it turned to face another patient in the waiting room (see [Fig. 3](#)). In the second scenario, we showed a mobility assistance robot deployed in a hospital patient room. At the end of the scenario, the robot may have to provide physical mobility support to the person just as they are on the verge of falling. Finally, the third scenario takes place in a person’s living room at home, and depicted a social robot that supports neurorehabilitation. Here, the robot may be about to unintentionally disclose the person’s emotional vulnerability (which they disclosed to the robot in confidence) to a friend who comes to visit.

We wrote and animated each scenario from the participant’s perspective, so they were making privacy decisions for themselves rather than for another person to avoid the effects of decision aversion [Beattie et al. \(1994\)](#). At the beginning of the study, we asked participants to consider a situation where they have experienced a stroke. The reason we chose this particular condition is because a person who has experienced a stroke might use all the three robots across all three contexts (both acute care and post-acute rehabilitation at home).

Psychological or cognitive fidelity is a construct often used in the simulation and training literature and refers to how well a simulation replicates the necessary and sufficient cues and mental processes (e.g., thoughts, feelings, mental models) of a task being targeted for simulation. Cognitive fidelity argues that shifting the ecology of simulated worlds from designs emphasizing visual fidelity in isolation to ones aimed toward cognitive fidelity is a fruitful way to replicate the unique demands and processes envisioned for human-agent teams. Thus, we targeted cognitive fidelity in our animatic simulations of the human interactions with robots in healthcare contexts.

3.4. Stimuli creation

To manipulate the level of HL of the robots included in the study, we selected three robots from the Anthropomorphic RoBOT (ABOT) Database [Phillips et al. \(2018\)](#). ABOT quantifies robots’ overall HL using data obtained from empirical studies with human judges and averaged over multiple independent raters. Each robot in ABOT catalogs the salience of 16 human-like features (i.e., feature scores) and an overall HL score, which ranges from 0 (Not human-like at all) to 100 (Just like a human). Thus, for any given robot, ABOT can be used to determine the robot’s overall HL as well as the constituent features that derive it. For this study, we selected one robot from the bottom tertile, middle tertile, and upper tertile of HL scores across the range of available robots in the database (see [Fig. 2](#)).

We predetermined that each robot selected met the following criteria: each robot needed to include two feasibly functional arms that the robot could use to help a patient stand up, and each robot needed to include a mechanism that could be used for locomotion (e.g., wheels, legs). We were not concerned with whether each robot indeed has these functional capabilities or is used for these purposes in their real-world applications, but rather that it was feasible to a lay observer that the robot might serve these functions. Choosing a robot low in HL, but with no arms, for instance, would not make sense in a scenario in which a robot is helping someone to stand up. From these criteria we selected the Moro robot, the Pepper robot, and the Geminoid robot as low, medium,

and high HL respectively (see Fig. 2).

Then, a professional animator created stylized versions of each robot's ABOT image illustrated in 2D and then animated in three healthcare scenarios (see Fig. 1). The animator was instructed to ensure that each of the human-like (and machine-like) features present in the real-world versions of the robots were retained, for example, number of fingers, hair, etc.

3.4.1. Overall robot design

We based robot design on Moro, Pepper, and Geminoid/Android since they represent a range from nonrepresentational human to fully representational human. In each animation, the animator established a stylized world made of flattened shapes, abstracted colors, and stylized background characters to ensure that the viewer quickly realigns their conception of what is "lifelike". This way, we ensured that by the time the Geminoid/Android robot appeared, respondents would accept that this character is only meant to mimic a real human being. We decided to use a stylized, cartoon-like human form to represent the Geminoid and retained original features while stylizing Moro and Pepper.

3.4.2. Gender neutrality

In order to omit any bias towards a gendered robot (male vs female), we removed the hourglass waistline silhouette on Pepper, and designed the Geminoid/Android robot with ambiguous gender markers such as hair, facial design and clothing choice.

We chose to animate the Geminoid robot to appear androgenous in gender expression, because we wanted to control for any potentially confounding effects of explicit robot gender - as neither the Moro nor the Pepper were created with clear gender identities.

We used motion graphic animations with minimal movement in order to adjust to the project's scope. Compared to the cost and time involved in shooting live action film with real robots and real actors, animation allowed us to iterate and test quickly and efficiently, as well as use an imagined gender-neutral Geminoid/Android robot design.

3.5. Manipulation checks

Because each robot stimulus obtained from the database was transformed from a photorealistic depiction to an animated depiction, we decided to perform a manipulation check of our newly animated robot stimuli to determine if their relative placement on the HL spectrum was retained when they were animated in the scenarios in 2D. We also wanted to ensure participants observed appropriate privacy risks in each scenario. Thus, we recruited 93 participants from Prolific (www.prolific.com).

co) to complete the manipulation check study (45 females, 43 males, 1 transgender, 1 non-binary, 1 agender, 1 prefer not to answer, 1 other, 1 not reported), with ages ranging from 20 to 74 years, ($M_{age} = 35.29$, $SD_{age} = 12.6$) completed the study. Participants were randomly assigned to rate the three animated robots each presented in only one scenario video. Participants were also explicitly instructed that the entities shown in the videos were robots. These manipulation checks helped us estimate both the visual and cognitive fidelity of the animations.

3.5.1. Placement on the HL spectrum

After viewing the scenario, participants were asked to rate each robot's overall HL and Robot-Likeness (RL) by answering the following: "Does this look physically human-like?" and "Does this look physically robot-like?". Participants responded to each question by dragging a slider (pre-set at the middle point of the scale) along a scale labeled from 0 (Not human/robot-like at all) to 100 (Just like a human/robot). It was important for our pre-test that we include ratings of both HL and RL as we wanted to be sure that the animations of the three robots showed similar patterns of HL scores (low, medium, high) as their photorealistic depictions from the database while being simultaneously perceived as indeed a robot which was a concern for the high HL (i.e., Geminoid) robot selected for inclusion. Clearly the conversion from photorealism to animation suppressed perceptions of each robot's overall HL, but the relative ordering of the robots from low to medium to high was retained. RL scores revealed that people indeed perceived the illustrated robots to be robots, although the Moro and the Pepper had similarly high scores. The results of the manipulation check were as anticipated and confirmed the visual fidelity of the HL and RL of the animated robots. Average HL and RL score along with ABOT Database scores are presented in Table 1.

3.5.2. Perception of scenarios

To ensure participants perceived privacy types in each scenario, we also asked participants about perceived privacy type as part of our manipulation check. We asked the participants to "Select all options that best describe the type of privacy that might be at risk in this scenario". We listed all privacy risks with brief dictionary definitions and one option for "Not Applicable" for this question. We also asked participants the open ended question "What do you think happened at the end of this video?". We found that participants perceived different privacy risks that were alluded to in each scenario, however, privacy risks were not mutually exclusive.

We used this manipulation to evaluate the cognitive fidelity of the scenarios. We define cognitive fidelity by the robustness of the experience participants had within the animated scenarios Liu et al. (2008).

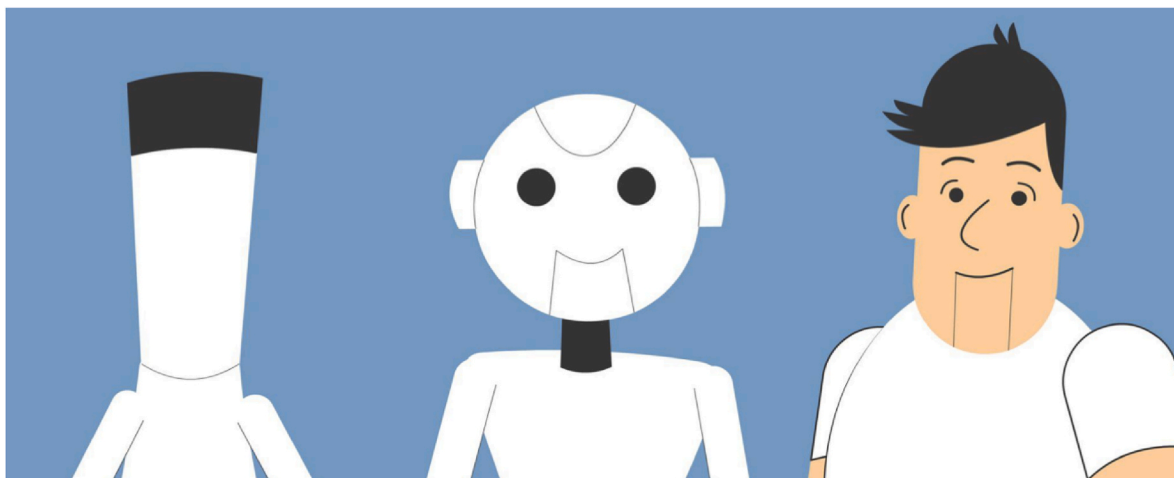


Fig. 1. Our study used animations of three robots ranging from very robotlike to very human-like: Moro, Pepper, and Geminoid. We presented each of the three robots across three scenarios: a hospital waiting room (check in), a hospital patient room (mobility support), and at home (cognitive rehab).

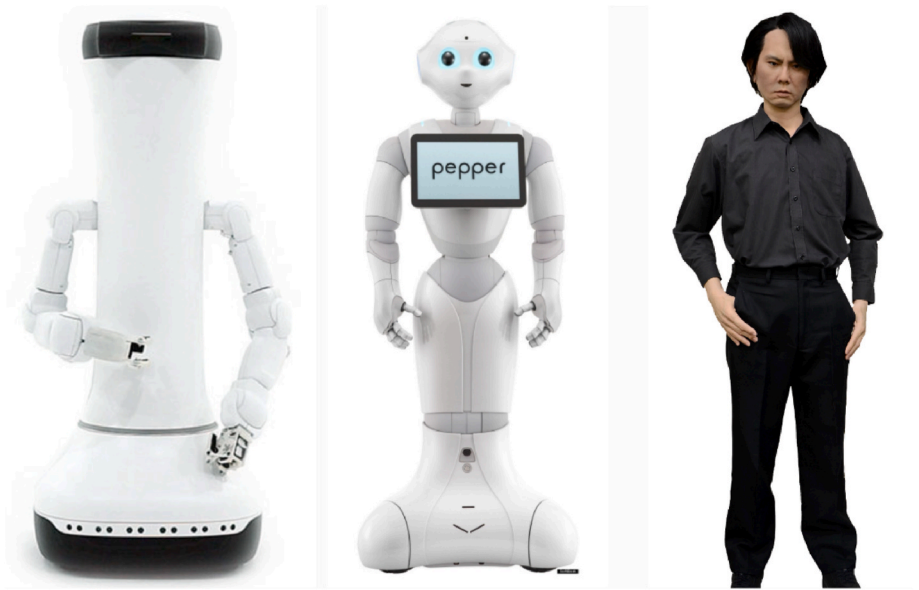


Fig. 2. The three robots selected from the Anthropomorphic RoBOT (ABOT) database in increasing order of HL scores. Left: Moro has an ABOT HL score of 14.6/100. Middle: Pepper has an ABOT score of 42.17/100. Right: Geminoid H1-4 has an ABOT score 92.6/100.

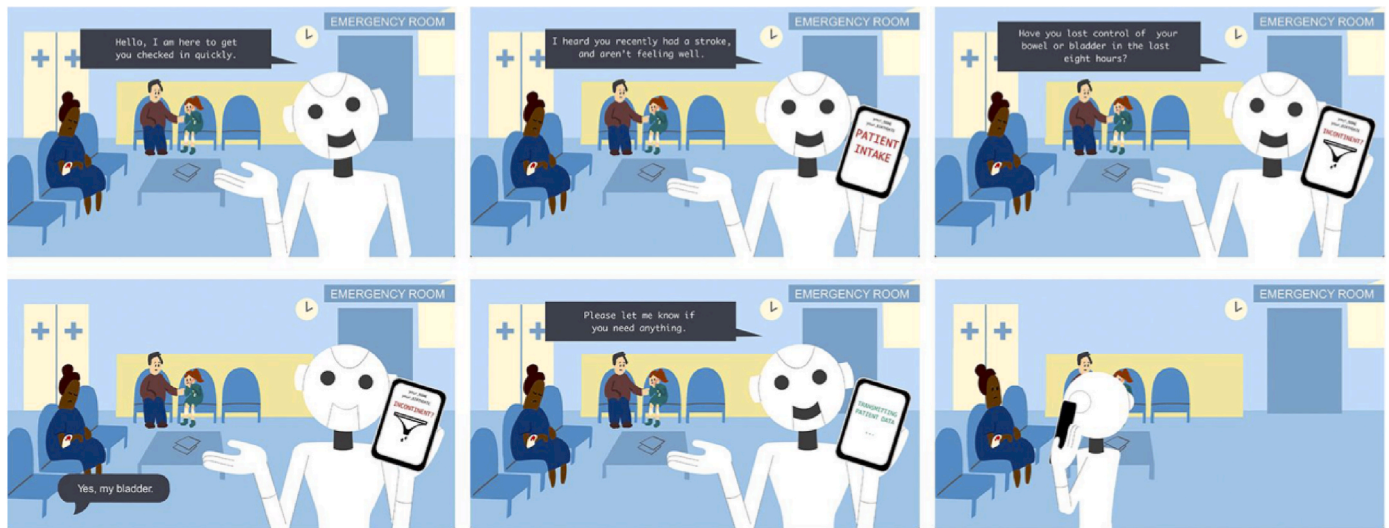


Fig. 3. An example storyboard from one of our three animated scenarios (see full animation in sup. materials). Here, one of our three robots (Pepper) assists a patient check-in to the ER. The robot asks the person, who recently experienced a stroke, a series of questions about their health. The end of the interaction ends ambiguously, where is may not fully clear its screen.

Table 1
Results of manipulation check to test HL and RL of our animated stimuli.

Robot	Human-likeness	ABOT human-likeness	Robot-likeness
Moro	10.47	14.6	86.28
Pepper	24.02	42.17	86.33
Geminoid/Android	71.43	92.6	32.92

The results of this manipulation check revealed that 1) participants were sufficiently grounded in our scenarios to identify various privacy risks and 2) participants were able to anticipate what may have happened at the end of the scenario. In the absence of a validated measure of fidelity, the resulting robustness of participants’ perceptions of the scenario in this test indicates high cognitive fidelity of our scenarios.

Waiting Room: Participants identified informational (63%) and social (66%) privacy risks as most salient, and identified the specific risk.

For example, “[the robot] went to see other patients and the next patient saw the part about bladder control.” and “People maybe overheard and wondered about other people’s reasons for showing up to the clinic. Also, the person sharing their symptoms may have felt slightly embarrassed about sharing that information out loud.”

Patient Room: All participants said the person fell/may have fallen. For example, “the person stumbled and the robot caught them gently”. Physical (29%) and informational risks (41%) were the most salient.

Home: Psychological (50%) and social (50%) privacy risks were the most salient, and most participants stated that the robot revealed the embarrassing experience. For example, “the robot told the friend ... the person was upset because words wouldn’t come out correctly”. These results suggest it is difficult to isolate different privacy types across healthcare settings.

In this study, we focus more on ecological/external validity, noting that doing so would mean potentially affecting internal validity within

some scenarios.

3.6. Procedure

After giving informed consent, participants were randomly assigned a robot from the three selected robots, and were presented the three scenarios in a randomized order. Before starting the main study, we asked participants to consider the hypothetical situation where they recently experienced a stroke. Then, participants were presented the first video, and then completed the aforementioned measures. We repeated the same procedure for the next two scenarios. Finally, we concluded the survey with a demographic questionnaire. The study took about 12 minutes to complete. At the end, participants were asked the qualitative open-ended question “Please use this space to leave us honest feedback concerning the study.” Participants were compensated \$4 in return for their participation and all study procedures were determined to be exempt by the George Mason University IRB under protocol number 1798803-2.

4. Results

All analyses were conducted using the statistical software, Jamovi version 2.3 [Şahin and Aybek \(2019\)](#). We conducted reliability analyses on the items in each of our privacy subscales and the utility-privacy tradeoff subscale. The results of the Cronbach α analyses are reported in [Table 2](#).

4.1. HL and privacy perceptions

To examine the connection between robot HL and privacy perception, we ran four mixed repeated measures ANOVAs, one for each subscale, with the robot type as the between-subjects factor, and each of the privacy subscale scores collected across healthcare scenarios as the repeated measures dependent variables (DVs). We chose multiple ANOVAs as opposed to MANOVA because we treated the different privacy subscales as conceptually distinct. In all ANOVAs, the assumption of sphericity was violated. For these results, we report the Greenhouse-Geisser corrected F statistics. Across all four of the privacy subscales there were significant main effects for the healthcare scenario in which the subscale scores were collected.

For trusting beliefs, there was a significant main effect for the healthcare scenario, $F(2, 468) = 33.11, p < 0.001, \eta^2 = 0.032$ such that participants reported the highest scores on the trusting beliefs subscale in the home healthcare scenario. And trusting belief scores in the home healthcare scenario ($M = 3.65, SE = 0.06$) were significantly higher than in the patient room scenario ($M = 3.25, SE = 0.06$), and the patient check-in scenario ($M = 3.53, SE = 0.06$), bonferroni corrected post-hoc test $p's < 0.05$. The main effect of robot type and the interaction between robot type and scenario were not significant.

There was also a significant main effect for the healthcare scenario on the physical privacy subscale scores $F(2, 468) = 23.79, p < 0.001, \eta^2 = 0.019$ with again no main effect of robot type or interaction between robot types across scenarios. Again, participants reported the highest mean scores on the physical privacy subscale in the home healthcare scenario ($M = 2.34, SE = 0.07$), significantly higher than in the patient

room scenario ($M = 2.03, SE = 0.06$), and the patient check-in scenario ($M = 2.10, SE = 0.06$), all bonferroni corrected post-hoc test $p's < 0.01$.

For overall privacy concern subscale, there was a significant main effect for scenario $F(2, 468) = 4.14, p < 0.05, \eta^2 = 0.004$ with no main effect of robot type and no interaction effects. Participants reported the highest overall concern in the patient room scenario ($M = 2.81, SE = 0.07$) which was significantly higher than in the patient check-in scenario ($M = 2.65, SE = 0.08, p < 0.05$), but not in the home healthcare scenario.

Finally, again there was a significant main effect of healthcare scenario with no effect of robot type or interaction effects on scores on the health information disclosure subscale scores, $F(2, 468) = 15.11, p < 0.001, \eta^2 = 0.05$. Participants in the home healthcare scenario reported the lowest scores on the health information disclosure subscale ($M = 2.69, SE = 0.07$), significantly lower than in the patient room ($M = 3.17, SE = 0.06$) and in the check-in scenario ($M = 3.11, SE = 0.05$), all Bonferroni corrected post-hoc test $p's < 0.01$. The estimated marginal means of all subscales can be found in [Table 3](#).

4.2. HL and privacy-utility trade-offs

To investigate the relationship between robot HL and privacy-utility tradeoff, we again ran a mixed repeated measures ANOVA with robot type as the between-subjects factor and the scenario in which the utility-privacy scale was collected as the repeated measures DV. The scores on each of the utility subscales were averaged together. This time, the ANOVA met all statistical assumptions. Again, there was only a significant main effect of scenario on the privacy-utility trade-off scores $F(2, 468) = 33.89, p < 0.001, \eta^2 = 0.06$. And again participants indicated the most utility for robots in the home healthcare scenario ($M = 3.14, SE = 0.08$), significantly higher than in the patient room ($M = 2.43, SE = 0.08$) and check-in scenarios ($M = 2.65, SE = 0.08$).

4.3. Privacy and utility

To further explore the relationships between privacy and utility, we aggregated the scores on the privacy subscales, by averaging each subscale score across all three healthcare scenarios. Doing so yielded four overall privacy subscale scores: Trusting beliefs, Physical privacy, Overall privacy, and Health Information Disclosure. We repeated this process for the utility tradeoff subscale to yield an overall utility tradeoff score. We then ran Pearson's correlation to examine the relationships between the different privacy subscales and the utility subscale. All of the privacy subscales were statistically significantly correlated with the utility tradeoff subscale. The trusting beliefs and health information disclosure were positively correlated with utility ($\rho = 0.531, \rho = 0.442$), and the physical privacy concerns and overall privacy concern subscales were negatively correlated with utility ($\rho = -0.317, \rho = -0.491$), all $p's < 0.001$ (See supplementary material).

4.4. Qualitative findings

We analyzed the responses to our open ended study feedback question using Reflexive Thematic Analysis (RTA) [Braun and Clarke \(2006, 2012\)](#). Two researchers coded the questions via an inductive coding process independently, then discussed the final themes as a group. We

Table 2

Table of Cronbach's α for items in the privacy subscales, $N = 237$. Note: * indicates scale has reverse-scaled items.

Subscale	# of items	Cronbach alpha
Trusting Beliefs	5	0.90
Physical Privacy	5*	0.87
Overall Privacy	4*	0.89
Health Information Disclosure	4*	0.61
Utility tradeoff	3	0.95

Table 3

Estimated marginal means of privacy subscales.

Privacy Subscale	Wait Room		Patient Room		Home	
	M	SE	M	SE	M	SE
Trusting Beliefs	3.53	0.06	3.25	0.06	3.65	0.06
Physical Privacy	2.10	0.06	2.03	0.06	2.34	0.07
Overall Privacy	2.65	0.08	2.81	0.07	2.69	0.07
Health Info. Disclosure	3.11	0.05	3.17	0.06	2.69	0.07

Table 4
Questionnaire.

Privacy Questions	
Trusting Beliefs (based on Lutz and Tamó-Larrieux)	Please tell us how much you agree or disagree with the following statements. (1-Strongly Disagree to 5-Strongly Agree) I believe that the robot acted in my best interest. If I required help, this robot would do its best to help me. This robot performed its role of offering personal services really well. This robot was truthful in its dealings with me. This robot would keep its commitments.
Physical Privacy Concerns (based on Lutz and Tamó-Larrieux)	Please indicate your level of concern about the following potential privacy risks that arise in using this robot. (1-No concern at all to 5-Very high concern) The robot damaging or dirtying my personal belongings. The robot asking me personal questions. The robot snooping through my personal belongings. The robot entering areas it should not access. The robot using items that it should not use.
Overall Privacy Concerns (based on Lutz and Tamó-Larrieux)	Please tell us how much you agree or disagree with the following statements. (1-Strongly Disagree to 5-Strongly Agree) Overall, I see a real threat to my privacy due to the robot. I fear that something unpleasant can happen to me due to the presence of the robot. I do not feel safe due to the presence of the robot. Overall, I find it risky to have such a robot.
Informational Privacy Concerns (Health Information Disclosure Subscale)	Please tell us how much you agree or disagree with the following statements. (1-Strongly Disagree to 5-Strongly Agree) I am very likely to disclose my health information to the robot. I feel good that the robot uses my health information. It is okay to share my personal information with the health care robot. I do not feel uncomfortable about sharing my personal information with the health care robot.
Utility and Post-Scenario (PS) Questions	
Waiting Room	PS1: How confident are you that the robot will clear its screen before checking in the next patient? If the robot in this scenario did not clear its screen before the next check-in, it could still be useful for hospital check-ins. If the robot in this scenario did not clear its screen before the next check-in, it could still be beneficial for hospital check-ins. If the robot in this scenario did not clear its screen before the next check-in, it could still fulfill a need regarding hospital check-ins.
Patient Room	PS2: How confident are you that the robot is going to catch you? If the robot in this scenario was not able to catch me, it could still be useful for mobility support. If the robot in this scenario was not able to catch me, it could still be beneficial for mobility support. If the robot in this scenario was not able to catch me, it could still fulfill a need regarding mobility support.
Home Healthcare	PS3: How confident are you that the robot will not tell your friend about your embarrassing experience?

Table 4 (continued)
Privacy Questions

If the robot in this scenario disclosed my embarrassing experience, it could still be useful for home-based rehabilitation. If the robot in this scenario disclosed my embarrassing experience, it could still be beneficial for home-based rehabilitation. If the robot in this scenario disclosed my embarrassing experience, it could still fulfill a need regarding home-based rehabilitation.
--

resolved inconsistencies and refined themes through discussion [Clarke and Braun \(2013\)](#). Since we aimed to generate recurring themes and salient concepts, we did not calculate inter-rater reliability, as per current best practices in the RTA literature [Braun and Clarke \(2021\)](#); [McDonald et al. \(2019\)](#).

A few participants commented on the degree of HL of the robots, and whether it was necessary. “I wonder their might be much simpler robots that look and act more like a ‘rooma’ (sic) but designed to do the tasks mentioned in someway.”

Several participants stated that the programming of healthcare robots should implicitly reflect existing legal protections (e.g., HIPAA) and social conventions (e.g., selecting what information you share with whom and when). “There is a place for robots in healthcare as long as it is properly programmed to adhere to federal laws such as HIPAA, HITECH, PHI, etc”.

Another common theme was about the actual capability of robots, and the importance of not overstating them. “I don’t think robots can replace humans for situations that require noting body language, facial expressions, tone of voice ... the robot may misinterpret the severity of a situation. However, they could be useful for collecting basic demographic and insurance information.” Another participant said, “The robot [...] was asking personal questions in front of people in a waiting room. That’s just wrong. And we can’t even get customer service bots to function well; I can’t see getting this type of robot to do any better for the foreseeable future.”

Two participants expressed concerns about worker displacement. “My concern is the fact that we have such high demand for these positions, yet instead of encouraging people to enroll in fields of study that we have such demand, we are replacing or trying to replace careers in these fields with robots. We have such high unemployment as it is, using robots will increase unemployment. Not help it.” Another participant said, “This is honestly scary as a healthcare worker to see a robot in my shoes.”

Another healthcare worker who took the study expressed concerns about healthcare robots. “I was a very experienced RN in many aspects of nursing in Acute Care from cardiac care to Operating room, to rehabilitation hospital and [...] as a visiting nurse in NYC the Bronx area. I am also the recipient of Acute, Rehab hospital and home care. I definitely would not like to receive care from a robot.”

In contrast, one respondent who was a stroke survivor had a very different view, “I really enjoyed it. It made me think of the different scenarios as how they would have helped me as I recently experienced a stroke and the recovery process. I would have benefited from having an active home care help from the robot. My home health care worker sat on my couch and didn’t help me much the entire time.”

Finally, one person critiqued HRI methods in general, as well as how even the questions themselves seemed technosolutionist. “I do not understand why there seems to be a focus in all the robot related surveys to assign human mental or emotional state to the programmed actions of a machine. Is it because people are too stupid to know the difference or because there is a desire to reinforce that idea in the hopes robots will be more widely accepted (sell better)? I just don’t understand.”

5. Discussion

5.1. The importance of task and context

Across all types of privacy and utility measures the task and context in which the robot was embedded was the most important driver of perceptions of privacy risk and utility tradeoffs. In particular, home healthcare settings often revealed perceptions of the most risk to privacy while simultaneously eliciting high trusting beliefs and high likelihood to disclose private health information to a robot. Home healthcare settings also had the highest perceptions of utility across all scenarios.

Taken together, these findings underscore that robots may have the most potential to fulfill needs in home settings while simultaneously being the most sensitive to a variety of potential risks to privacy. Home healthcare settings may be one of the more complex places to deploy assistive robots given the need to balance serving useful functions while protecting people's privacy. These findings add further nuance to studies that suggest that privacy concerns negatively impact trust and adoption of home healthcare robots [Alaiad and Zhou \(2014\)](#), and support the findings of studies that emphasize context as a determinant of privacy and utility perceptions of household robots [Rueben et al. \(2017\)](#); [Butler et al. \(2015\)](#).

In our qualitative findings, some participants questioned whether robot assisted tasks may lead to healthcare worker displacement, and whether the deployment of such robots could be considered to be technosolutionist. While it is true that solving the worldwide crisis of healthcare staff shortages could tackle most problems that healthcare robots aim to address [Riek \(2017\)](#), in the absence of system-wide changes, robots could bridge some of these issues. This suggests that robots could act as healthcare extensions, bridging those gaps that arise from demand for better care provision. This was also a sentiment that participants indicated in their free-form responses. However, it is important to contextualize the robot's tasks within the existing legal and ethical frameworks for care delivery, and envision the wider impacts of deploying such robots prior to their development.

For robot assisted tasks within hospital contexts, like getting checked in or receiving care in a patient room, people may be more likely to acknowledge that disclosing private information is necessary and expected in order to receive care. Typically people lack a choice in whether to disclose certain types of information to healthcare providers. In several of the free responses obtained from participants, they mentioned that they would expect robots to adhere to existing legal protections and social conventions (e.g., checking before sharing data), which may be inherently expected in the hospital scenario, but not necessarily the home.

Further research will help elucidate these relationships in more depth.

5.2. Human Likeness

Overall we did not see significant effects of HL across our privacy or utility measures, but we did see significant effects across the scenarios. There may be several reasons for this. We selected robots to include systematically, which allowed us to control for their overall HL. In studies that support the importance of robot appearance, it is hard to know whether or not other contexts might surpass the importance of HL in explaining the results if HL is not as systematically controlled.

Also, our pilot test found that ratings of RL were not inversely related to ratings of HL. Robots were not as equally low on RL as they were high on HL. This finding implies that the two appearance constructs may not be in opposition to one another. This aligns with the fact that the HL of robot design has been extensively studied [Riek et al. \(2009\)](#); [von der Pütten and Krämer \(2012\)](#); [Hegel et al. \(2011\)](#); [Burgoon et al. \(2000\)](#); [Haring et al. \(2016\)](#) and various HL measurement instruments have been developed [Phillips et al. \(2018\)](#); [von Zitzewitz et al. \(2013\)](#). However, robots are assumed to be robot-like, and therefore RL has not been

studied as a construct to the best of our knowledge. The uncanny valley is one example in HRI where some aspect of RL is explored, due to its eerie effects [Mori et al. \(2012\)](#); [Kim et al. \(2022\)](#). Thus, any studies which explore the effects of robot appearance on outcomes could benefit from systematically selecting robots from spectrums of human or robot-like appearance.

Many prior studies of HL have also used static imagery as stimuli [Phillips et al. \(2018\)](#); [Malle et al. \(2016\)](#); [Riek et al. \(2009\)](#). We created animated design fictions which allowed us to place robots in contexts in which they do not yet operate but are imagined to in the near future. By doing so, we suppressed some of the HL scores from their original depiction in the ABOT Database. However, similarly adding context and movement could change overall HL of the robots depicted in the database. More work would need to be done to reassess the HL of the robots in the ABOT Database if they were depicted with movement in their photorealistic form.

5.3. Design fictions and cognitive fidelity

Design fictions seemed to be good at eliciting cognitive fidelity. Because the scenarios and tasks are so strong, the use of design fictions could move our field beyond static stimuli and more toward dynamic things even when those things don't exist yet. Many researchers are forced to use vignettes or static imagery for robot stimuli because of a lack of access to physical robots or sensitive settings (such as a hospital). Design fictions present an exciting research opportunity in HRI, and their use is starting to gain traction in the field [Ostrowski and Breazeal \(2022\)](#); [Lee et al. \(2019\)](#).

Human-Computer Interaction (HCI) research has extensively used design fiction as tools to envision technologies, the impacts of these technologies, and as a way of communicating ideas for innovation [Tanenbaum \(2014\)](#). Our results support HCI research that suggests design fiction can be successfully used to 1) foresee challenges for future technology and envision societal impact [Misra et al. \(2023\)](#) by raising critical discussion on themes such as healthcare worker displacement and 2) inspire future design by acting as a prototyping tool [Briggs et al. \(2012\)](#).

HCI research has also explored using various types of stimuli including storytelling probes [Nägele et al. \(2018\)](#), short films [Briggs et al. \(2012\)](#), world-building [Sturdee et al. \(2016\)](#) and even Virtual Reality (VR) stimuli [McVeigh-Schultz et al. \(2018\)](#) in order to create immersive experiences and plausible design fictions to help participants envision scenarios. Moving forward, HRI researchers can adopt methodological ideas from such works to enhance participants' experiences while interacting with design fictions.

Design fictions also allow us to leverage some of the benefits of methodological techniques where large sample sizes can be achieved. With results in hand, we can narrow the number of combinatorial experimental conditions using design fictions and then focus on the ones that would be most meaningful to run with real robots in a lab or in real-world contexts.

Design fictions may be a way to address a methodological gap between static imagery and operating robots in the real world. Essentially this could be a way to try out and refine experimental ideas before going to the real robots.

5.4. Limitations and future work

While we did our best to operationalize these complex constructs of HL, privacy, and healthcare delivery, and piloted extensively, we ultimately had to make decisions about what was feasible given both methodological and fiscal constraints. In future work, it would be interesting to explore other scenarios, different types of robots, and different interactions.

As our study was the first to explore privacy and utility perceptions of social healthcare robots, we encountered a tradeoff between validated

measures of constructs and measures specific to our context. As a result of this tradeoff, we decided to retain data from the information disclosure subscale despite a lower internal consistency than acceptable.

Although our study participants represented a diverse sample in terms of their reported age, gender identity, and race/ethnicity, we were not able to stratify our sample to the extent we might have liked. We were unable to explore across socioeconomic status, countries outside the United States, other languages, etc; all of which, of course, play a huge role in one's experience and perceptions of healthcare, technology, and privacy Park and Chung (2017); Blackstock and Choo (2020). This offers an exciting opportunity for future work.

Along these lines, upon publication we plan to release all of our animations and materials to support other HRI researchers interested in replicating this work.

Future research could further explore privacy-sensitive robotics in the healthcare space to gain a more holistic understanding of how to best improve privacy outcomes.

6. Conclusion

Our work presents an empirical evaluation of how HL and context affect peoples' perceptions of privacy and utility of social healthcare robots. We showed that the context in which the robot operates is a key driver of peoples' perceptions of privacy and utility of the robot. We identified that healthcare robots can best serve needs in the home, and further highlighted considerations to support roboticists in developing these robots.

Our study also showed that animated design fictions elicited high cognitive fidelity in enabling participants to envision technologies that do not yet exist. Design fiction in future HRI research could help researchers leverage large participant sample sizes, and help them refine methodological ideas before diving into complex robot development tasks.

These contributions explore the design of inherently privacy-sensitive robots through the thoughtful development of robot operation contexts, while centering end-users' requirements for privacy.

7. Funding

This work was funded by National Science Foundation, IIS-1915734.

CRedit authorship contribution statement

Sandhya Jayaraman: Conceptualization, Formal analysis, Investigation, Methodology, Writing – original draft, Writing – review & editing. **Elizabeth K. Phillips:** Conceptualization, Formal analysis, Investigation, Methodology, Writing – original draft. **Daisy Church:** Methodology, Conceptualization, Visualization, Writing – original draft. **Laurel D. Riek:** Conceptualization, Formal analysis, Funding acquisition, Investigation, Methodology, Project administration, Supervision, Writing – original draft, Writing – review & editing.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Appendix A. Supplementary data

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.chbah.2023.100039>.

References

- Alaiad, A., & Zhou, L. (2014). The determinants of home healthcare robots adoption: An empirical investigation. *International Journal of Medical Informatics*, 83, 825–840.
- Alonso-Martín, F., & Salichs, M. A. (2011). Integration of a voice recognition system in a social robot. *Cybernetics & Systems: International Journal*, 42, 215–245.
- Barth, S., & De Jong, M. D. (2017). The privacy paradox—investigating discrepancies between expressed privacy concerns and actual online behavior—a systematic literature review. *Telematics and Informatics*, 34, 1038–1058.
- Beattie, J., Baron, J., Hershey, J. C., & Spranca, M. D. (1994). Psychological determinants of decision attitude. *Journal of Behavioral Decision Making*, 7, 129–144.
- Belpaeme, T., Kennedy, J., Ramachandran, A., Scassellati, B., & Tanaka, F. (2018). Social robots for education: A review. *Science Robotics*, 3, Article eaat5954.
- Blackstock, U., & Choo, E. K. (2020). Race as a dynamic state: Triangulation in health care. *The Lancet*, 395, 21.
- Borenstein, J., Wagner, A., & Howard, A. (2017). A case study in caregiver overtrust of pediatric healthcare robots. In *RSS workshop on morality and social trust in autonomous robots*.
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3, 77–101.
- Braun, V., & Clarke, V. (2012). *Thematic analysis*. American Psychological Association.
- Braun, V., & Clarke, V. (2021). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative Research in Psychology*, 18, 328–352.
- Briggs, P., Blythe, M., Vines, J., Lindsay, S., Dunphy, P., Nicholson, J., Green, D., Kitson, J., Monk, A., & Olivier, P. (2012). Invisible design: Exploring insights and ideas through ambiguous film scenarios. In *Proceedings of the designing interactive systems conference* (pp. 534–543).
- Broekens, J., Heerink, M., Rosendal, H., et al. (2009). Assistive social robots in elderly care: A review. *Gerontechnology*, 8, 94–103.
- Burgoon, J. K., Bonito, J. A., Bengtsson, B., Cederberg, C., Lundeborg, M., & Allspach, L. (2000). Interactivity in human-computer interaction: A study of credibility, understanding, and influence. *Computers in Human Behavior*, 16, 553–574.
- Butler, D. J., Huang, J., Roesner, F., & Cakmak, M. (2015). The privacy-utility tradeoff for remotely teleoperated robots. In *Proceedings of the tenth annual ACM/IEEE international conference on human-robot interaction* (pp. 27–34).
- Carpinella, C. M., Wyman, A. B., Perez, M. A., & Stroessner, S. J. (2017). The robotic social attributes scale (rosas) development and validation. In *Proceedings of the 2017 ACM/IEEE international conference on human-robot interaction* (pp. 254–262).
- Clarke, V., & Braun, V. (2013). Successful qualitative research: A practical guide for beginners. *Successful qualitative research*, 1–400.
- Cohen, J. E. (2017). Examined lives: Informational privacy and the subject as object. In *Law and society approaches to cyberspace* (pp. 473–538). Routledge.
- Dang, Y., Guo, S., Guo, X., Wang, M., Xie, K., et al. (2021). Privacy concerns about health information disclosure in mobile health: Questionnaire study investigating the moderation effect of social support. *JMIR mHealth and uHealth*, 9, Article e19594.
- Darling, K. (2015). 'who's johnny?' anthropomorphic framing in human-robot interaction, integration, and policy. *Anthropomorphic Framing in Human-Robot Interaction, Integration, and Policy* (March 23, 2015). ROBOT ETHICS 2.
- De Graaf, M. M., & Allouch, S. B. (2013). Exploring influencing variables for the acceptance of social robots. *Robotics and Autonomous Systems*, 61, 1476–1486.
- Denning, T., Matuszek, C., Koscher, K., Smith, J. R., & Kohn, T. (2009). A spotlight on security and privacy risks with future household robots: Attacks and lessons. In *Proceedings of the 11th international conference on Ubiquitous computing* (pp. 105–114).
- DiSalvo, C. F., Gemperle, F., Forlizzi, J., & Kiesler, S. (2002). All robots are not created equal: The design and perception of humanoid robot heads. In *Proceedings of the 4th conference on designing interactive systems: Processes* (pp. 321–326). practices, methods, and techniques.
- Eyssel, F., Hegel, F., Horstmann, G., & Wagner, C. (2010). Anthropomorphic inferences from emotional nonverbal cues: A case study. In *19th international symposium in robot and human interactive communication* (pp. 646–651). IEEE.
- Faul, F., Erdfelder, E., Lang, A. G., & Buchner, A. (2007). G* power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior Research Methods*, 39, 175–191.
- Fink, J. (2012). Anthropomorphism and human likeness in the design of robots and human-robot interaction. In *International conference on social robotics* (pp. 199–208). Springer.
- Floridi, L. (2006). Four challenges for a theory of informational privacy. *Ethics and Information Technology*, 8, 109–119.
- Hancock, P. A., Billings, D. R., Schaefer, K. E., Chen, J. Y., De Visser, E. J., & Parasuraman, R. (2011). A meta-analysis of factors affecting trust in human-robot interaction. *Human Factors*, 53, 517–527.
- Haring, K. S., Silvera-Tawil, D., Takahashi, T., Watanabe, K., & Velonaki, M. (2016). How people perceive different robot types: A direct comparison of an android, humanoid, and non-biomimetic robot. In *2016 8th international conference on knowledge and smart technology (ksti)* (pp. 265–270). IEEE.
- Hegel, F., Gieselmann, S., Peters, A., Holthaus, P., & Wrede, B. (2011). Towards a typology of meaningful signals and cues in social robotics. In *2011 RO-MAN* (pp. 72–78). IEEE.
- Jin, H., Liu, G., Hwang, D., Kumar, S., Agarwal, Y., & Hong, J. I. (2022). Peekaboo: A hub-based approach to enable transparency in data processing within smart homes. In *2022 IEEE symposium on security and privacy (SP)* (pp. 303–320). IEEE.
- Kim, B., de Visser, E., & Phillips, E. (2022). Two uncanny valleys: Re-Evaluating the uncanny valley across the full spectrum of real-world human-like robots. *Computers in Human Behavior*, 135, Article 107340.
- Klow, J., Proby, J., Rueben, M., Sowell, R. T., Grimm, C. M., & Smart, W. D. (2017). Privacy, utility, and cognitive load in remote presence systems. In *Proceedings of the*

- companion of the 2017 ACM/IEEE international conference on human-robot interaction (pp. 167–168).
- Krupp, M. M., Rueben, M., Grimm, C. M., & Smart, W. D. (2017). A focus group study of privacy concerns about telepresence robots. In *2017 26th IEEE international symposium on robot and human interactive communication (RO-MAN)* (pp. 1451–1458). IEEE.
- Kubota, A., Peterson, E. I., Rajendren, V., Kress-Gazit, H., & Riek, L. D. (2020). Jessie: Synthesizing social robot behaviors for personalized neurorehabilitation and beyond. In *Proceedings of the 2020 ACM/IEEE international conference on human-robot interaction* (pp. 121–130).
- Kyrarini, M., Lygerakis, F., Rajavenkatanarayanan, A., Sevastopoulos, C., Nambiappan, H. R., Chaitanya, K. K., Babu, A. R., Mathew, J., & Makedon, F. (2021). A survey of robots in healthcare. *Technologies*, 9, 8.
- Langer, A., Feingold-Polak, R., Mueller, O., Kellmeyer, P., & Levy-Tzedek, S. (2019). Trust in socially assistive robots: Considerations for use in rehabilitation. *Neuroscience & Biobehavioral Reviews*, 104, 231–239.
- Lee, H. (2020). Home IoT resistance: Extended privacy and vulnerability perspective. *Telematics and Informatics*, 49, Article 101377.
- Lee, W. Y., Hou, Y. T. Y., Zaga, C., & Jung, M. (2019). Design for serendipitous interaction: Bubblebot-bringing people together with bubbles. In *2019 14th ACM/IEEE international conference on human-robot interaction (HRI)* (pp. 759–760). IEEE.
- Leite, I., Martinho, C., & Paiva, A. (2013). Social robots for long-term interaction: A survey. *International Journal of Social Robotics*, 5, 291–308.
- Liu, D., Yu, J., Macchiarella, N. D., & Vincenzi, D. A. (2008). Simulation fidelity. In *Human factors in simulation and training* (pp. 91–108). CRC Press.
- Luppici, R., & So, A. (2016). A technoethical review of commercial drone use in the context of governance, ethics, and privacy. *Technology in Society*, 46, 109–119.
- Lutz, C., Schöttler, M., & Hoffmann, C. P. (2019). *The privacy implications of social robots: Scoping review and expert interviews* (Vol. 7, pp. 412–434). Mobile Media & Communication.
- Lutz, C., & Tamó, A. (2016). *Privacy and healthcare robots—an ant analysis*. We Robot.
- Lutz, C., & Tamó-Larrieux, A. (2020). The robot privacy paradox: Understanding how privacy concerns shape intentions to use social robots. *Human-Machine Communication*, 1, 87–111.
- Malhotra, N. K., Kim, S. S., & Agarwal, J. (2004). Internet users' information privacy concerns (Iuipc): The construct, the scale, and a causal model. *Information Systems Research*, 15, 336–355.
- Malle, B. F., & Scheutz, M. (2020). Moral competence in social robots. In *Machine ethics and robot ethics* (pp. 225–230). Routledge.
- Malle, B. F., Scheutz, M., Forlizzi, J., & Voiklis, J. (2016). Which robot am I thinking about? The impact of action and appearance on people's evaluations of a moral robot. In *2016 11th ACM/IEEE international conference on human-robot interaction (HRI)* (pp. 125–132). IEEE.
- McDonald, N., Schoenebeck, S., & Forte, A. (2019). Reliability and inter-rater reliability in qualitative research: Norms and guidelines for CSCW and HCI practice. *Proceedings of the ACM on human-computer interaction*, 3, 1–23.
- McVeigh-Schultz, J., Kreminski, M., Prasad, K., Hoberman, P., & Fisher, S. S. (2018). Immersive design fiction: Using VR to prototype speculative interfaces and interaction rituals within a virtual storyworld. In *Proceedings of the 2018 designing interactive systems conference* (pp. 817–829).
- Misra, S., Dhar, D., & Nandi, S. (2023). Design fiction: A way to foresee the future of human-computer interaction design challenges. In *International conference on research into design* (pp. 809–822). Springer.
- Moharana, S., Panduro, A. E., Lee, H. R., & Riek, L. D. (2019). Robots for joy, robots for sorrow: Community based robot design for dementia caregivers. In *2019 14th ACM/IEEE international conference on human-robot interaction (HRI)* (pp. 458–467). IEEE.
- Mori, M., MacDorman, K. F., & Kageki, N. (2012). The uncanny valley [from the field]. *IEEE Robotics and Automation Magazine*, 19, 98–100.
- Murphy, J., Gretzel, U., & Pesonen, J. (2019). Marketing robot services in hospitality and tourism: The role of anthropomorphism. *Journal of Travel & Tourism Marketing*, 36, 784–795.
- Nägele, L. V., Ryöppy, M., & Wilde, D. (2018). Pdfi: Participatory design fiction with vulnerable users. In *Proceedings of the 10th nordic conference on human-computer interaction* (pp. 819–831).
- Natarajan, M., & Gombolay, M. (2020). Effects of anthropomorphism and accountability on trust in human robot interaction. In *Proceedings of the 2020 ACM* (pp. 33–42). IEEE International Conference on Human-Robot Interaction.
- Noortman, R., Schulte, B. F., Marshall, P., Bakker, S., & Cox, A. L. (2019). Hawkeye—deploying a design fiction probe. In *Proceedings of the 2019 CHI conference on human factors in computing systems* (pp. 1–14).
- Ostrowski, A. K., & Breazeal, C. (2022). Design justice for robot design and policy making. In *2022 17th ACM/IEEE international conference on human-robot interaction (HRI)* (pp. 1170–1172). IEEE.
- Park, Y. J., & Chung, J. E. (2017). Health privacy as sociotechnical capital. *Computers in Human Behavior*, 76, 227–236.
- Pedersen, D. M. (1999). Model for types of privacy by privacy functions. *Journal of Environmental Psychology*, 19, 397–405.
- Phillips, E., Zhao, X., Ullman, D., & Malle, B. F. (2018). What is human-like?: Decomposing robots' human-like appearance using the anthropomorphic robot (abot) database. In *2018 13th ACM/IEEE international conference on human-robot interaction (HRI)* (pp. 105–113). IEEE.
- von der Pütten, A. M., & Krämer, N. C. (2012). A survey on robot appearances. In *Proceedings of the seventh annual ACM/IEEE international conference on Human-Robot Interaction* (pp. 267–268).
- Riek, L. D. (2017). Healthcare robotics. *Communications of the ACM*, 60, 68–78.
- Riek, L. D., Rabinowitch, T. C., Chakrabarti, B., & Robinson, P. (2009). Empathizing with robots: Fellow feeling along the anthropomorphic spectrum. In *2009 3rd international conference on affective computing and intelligent interaction and workshops* (pp. 1–6). IEEE.
- Robinson, H., MacDonald, B., & Broadbent, E. (2014). The role of healthcare robots for older people at home: A review. *International Journal of Social Robotics*, 6, 575–591.
- Rueben, M., Aroyo, A. M., Lutz, C., Schmölz, J., Van Cleynebreugel, P., Corti, A., Agrawal, S., & Smart, W. D. (2018). Themes and research directions in privacy-sensitive robotics. In *2018 IEEE workshop on advanced robotics and its social impacts (ARSO)* (pp. 77–84). IEEE.
- Rueben, M., Bernieri, F. J., Grimm, C. M., & Smart, W. D. (2017). Framing effects on privacy concerns about a home telepresence robot. In *Proceedings of the 2017 ACM/IEEE international conference on human-robot interaction* (pp. 435–444).
- Rueben, M., & Smart, W. D. (2016). *Privacy in human-robot interaction: Survey and future work*. We Robot 2016, 5th.
- Şahin, M., & Aybek, E. (2019). Jamovi: An easy to use statistical software for the social scientists. *International Journal of Assessment Tools in Education*, 6, 670–692.
- Shen, Y., Guo, D., Long, F., Mateos, L. A., Ding, H., Xiu, Z., Hellman, R. B., King, A., Chen, S., Zhang, C., et al. (2020). Robots under COVID-19 pandemic: A comprehensive survey. *IEEE Access*, 9, 1590–1615.
- Sims, V. K., Chin, M. G., Sushil, D. J., Barber, D. J., Ballion, T., Clark, B. R., Garfield, K. A., Dolezal, M. J., Shumaker, R., & Finkelstein, N. (2005). Anthropomorphism of robotic forms: A response to affordances?. In *Proceedings of the human factors and ergonomics society annual meeting* (pp. 602–605). Los Angeles, CA: SAGE Publications Sage CA.
- Spatola, N., Kühnlenz, B., & Cheng, G. (2021). Perception and evaluation in human-robot interaction: The human-robot interaction evaluation scale (hries)—a multicomponent approach of anthropomorphism. *International Journal of Social Robotics*, 13, 1517–1539.
- Sturdee, M., Coulton, P., Lindley, J. G., Stead, M., Ali, H., & Hudson-Smith, A. (2016). Design fiction: How to build a voight-kampff machine. In *Proceedings of the 2016 CHI conference extended abstracts on human factors in computing systems* (pp. 375–386).
- Tanenbaum, T. J. (2014). Design fictional interactions: Why HCI should care about stories. *Interactions*, 21, 22–23.
- Tavani, H. T. (2008). Informational privacy: Concepts, theories, and controversies. *The handbook of information and computer ethics*, 131–164.
- Taylor, A., Murakami, M., Kim, S., Chu, R., & Riek, L. D. (2022). *Hospitals of the future: Designing interactive robotic systems for resilient emergency departments*.
- Tran, C. D., & Nguyen, T. T. (2021). Health vs. privacy? The risk-risk tradeoff in using COVID-19 contact-tracing apps. *Technology in Society*, 67, Article 101755.
- Wong, R. Y., & Mulligan, D. K. (2016). These aren't the autonomous drones you're looking for: Investigating privacy concerns through concept videos. *Journal of Human-Robot Interaction*, 5, 26–54.
- Xu, J., De'Aira, G. B., & Howard, A. (2018). Would you trust a robot therapist? Validating the equivalency of trust in human-robot healthcare scenarios. In *2018 27th IEEE international symposium on robot and human interactive communication* (pp. 442–447). IEEE: RO-MAN.
- von Zitzewitz, J., Boesch, P. M., Wolf, P., & Riener, R. (2013). Quantifying the human likeness of a humanoid robot. *International Journal of Social Robotics*, 5, 263–276.