

Mirror Skin: In Situ Visualization of Robot Touch Intent on Robot Skin

David Wagmann
Saarland University,
Saarland Informatics Campus
Saarbrücken, Germany
wagmann@cs.uni-saarland.de

Matti Krüger
Honda Research Institute Europe
Offenbach/Main, Germany
matti.krueger@honda-ri.de

Chao Wang
Honda Research Institute Europe
Offenbach/Main, Germany
chao.wang@honda-ri.de

Michael Gienger
Honda Research Institute Europe
Offenbach/Main, Germany
michael.gienger@honda-ri.de

Jürgen Steimle
Saarland University,
Saarland Informatics Campus
Saarbrücken, Germany
steimle@cs.uni-saarland.de

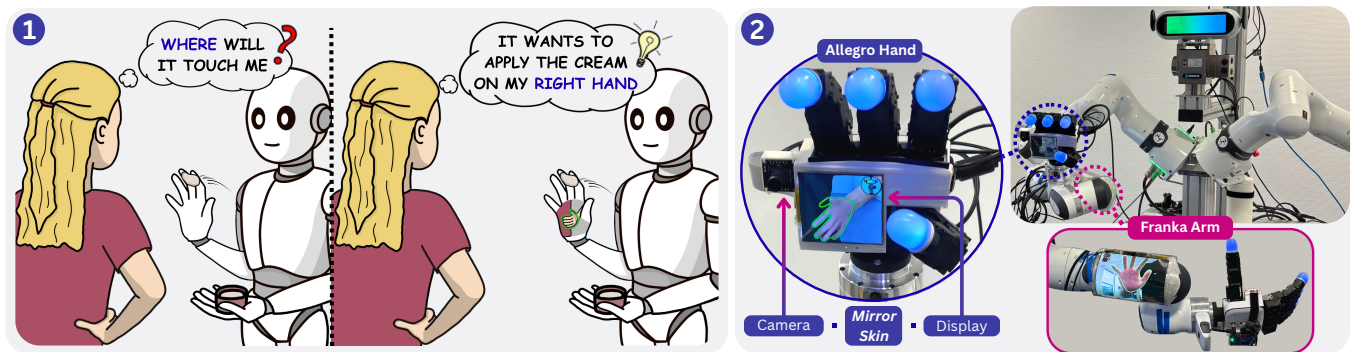


Figure 1: *Mirror Skin* communicates a robot's touch intent through mirror-like visual reflections on robot skin (1). We developed the *Mirror Skin* concept through a VR design exploration and built a proof-of-concept system (2). User studies in VR and with the physical prototype showed that *Mirror Skin* efficiently conveys robot touch intent and enhances the user experience.

Abstract

Effective communication of robot touch intent is essential for safe and predictable physical human-robot interaction. While intent communication has been widely studied, existing approaches lack the spatial specificity and semantic depth necessary to efficiently convey robot touch intent. We present *Mirror Skin*, a cephalopod-inspired concept that mirrors in-situ visual representations of a human's body parts onto the corresponding robot's touch region to communicate *who* shall initiate touch, *where* it will occur, and *when* it is imminent. We informed the design of *Mirror Skin* through a structured design exploration with experts and demonstrate the real-world feasibility of *Mirror Skin* with a proof-of-concept prototype. User studies in VR and with the physical prototype showed that *Mirror Skin* significantly improves accuracy and response times for interpreting touch intent and improves the user experience during physical human-robot interactions.



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

UIST '26, Detroit, MI, USA

© 2026 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-2856-3/2026/11

<https://doi.org/10.1145/3830398.3830547>

CCS Concepts

• Human-centered computing → Interactive systems and tools.

Keywords

Touch intent, robot skin, physical human-robot interaction

ACM Reference Format:

David Wagmann, Matti Krüger, Chao Wang, Michael Gienger, and Jürgen Steimle. 2026. Mirror Skin: In Situ Visualization of Robot Touch Intent on Robot Skin. In *The 39th Annual ACM Symposium on User Interface Software and Technology (UIST '26)*, November 02–05, 2026, Detroit, MI, USA. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3830398.3830547>

1 Introduction

A new generation of robots is emerging that engages in physical interaction with humans across various domains, including manufacturing [46, 60], healthcare [10, 23], and social interactions [64]. In such physical human-robot interactions (pHRI), it is important for humans to know when the robot intends to make contact [68], as this awareness is fundamental for creating predictability [10], ensuring safety [56], and improving task performance [74].

However, despite extensive research on the effects of robot-initiated touch in human-robot interaction (HRI), there remains a

gap in understanding how to effectively and efficiently communicate a robot's touch intent. While a substantial body of work addresses general robot intent communication, only few prior studies have included signaling of touch intent. These have predominantly relied on high-level verbal cues (notably, [10, 31]), which may not always be favorable [10] or effective [67], foregoing alternatives that may enhance clarity and user experience (UX). This is problematic, as uncertainty about *who* shall initiate touch and *where* exactly it will occur can lead to expectation mismatches, potentially lowering the efficiency and safety of the interaction. Therefore, we argue that it is important to explore dedicated communication concepts for human-robot touch to convey rich semantic feedback.

To address this gap, we draw inspiration from animals—such as cephalopods, lizards, and chameleons—that utilize visual color and texture changes as a powerful form of communication [36, 65, 69]. Cephalopods are particularly interesting due to their ability to alter skin color and texture with a high spatial and temporal resolution, for instance, to reflect parts of the environment onto their skin and to create visual highlights to direct attention toward, or away from, specific body regions [33]. These capabilities offer unique design opportunities for dynamic robot skin-based communication.

As a result, we present *Mirror Skin*: a cephalopod-inspired concept that employs high-resolution visual feedback on the robot's skin to communicate robot touch intent. Using the metaphor of a mirror, it communicates *when* a touch event is imminent, *who* shall initiate touch and *where* it will occur. This is achieved by mirroring live reflections of the interacting human body part onto the surface of the robot's touching body part (cf., Figure 1). For instance, mirroring the human's arm onto the robot's moving end effector can indicate the robot's intent to initiate contact at that body region. Conversely, reflecting the human's hand onto the robot's stationary arm can signal an invitation to touch the robot at that location.

We chose the mirror metaphor as a form of semantically meaningful visual feedback, since mirror reflections are a powerful mechanism for self-recognition [5]. Humans develop the innate ability to identify their own bodies and specific body parts in reflections from an early age [70], making it an efficient and intuitive visual mechanism for body-centered communication [41].

To systematically develop our concept, we conducted a broad design exploration in virtual reality (VR). We identified six key dimensions that shape the design of *Mirror Skin* and iteratively prototyped design variants to maximize visual clarity and salience. We then evaluated those generated design variants in an exploratory VR study ($N = 7$) with robotics and design experts and refined our design of *Mirror Skin*. Using our refined implementation, we conducted a lab study ($N = 12$) in VR to quantitatively and qualitatively assess its effectiveness for conveying touch intent before and during touch motions. The results demonstrate that *Mirror Skin* significantly improves the accuracy and response time for interpreting robot touch intent compared to a baseline condition consisting of robot gestures and gaze. Finally, we implemented a physical proof-of-concept system incorporating high-resolution displays and wide-angle cameras to instantiate *Mirror Skin* in a real-world setting, enabling evaluation of its effects on intent communication as well as on pragmatic and hedonic UX. A lab study ($N = 12$) confirmed that compared to robot gestures, *Mirror Skin* communicated the robot's touch intent more clearly and yielded better pragmatic

and hedonic UX, including higher confidence to engage in physical interaction with the robot.

In summary, this paper contributes:

- (1) *Mirror Skin*, a cephalopod-inspired concept for semantically rich visual communication of touch-related intentions on robot skin, including *who* shall initiate touch, *when* touch is imminent and *where* the touch is happening.
- (2) An iterative design exploration that introduces and investigates six key components of *Mirror Skin*, validated through an exploratory study with domain experts.
- (3) Findings from a VR user study that validates the suitability of *Mirror Skin* for efficiently conveying robot touch intent.
- (4) A physical proof-of-concept implementation of *Mirror Skin* and real-world user study demonstrating improved UX.

2 Related Work

This work is informed by prior work on pHRI and communication of robot intent in HCI and HRI.

2.1 Physical Human-Robot Interaction

Physical human-robot interaction has gained increasing attention due to the central role of touch in human social and cooperative interactions. Recent advances in robotic sensing technologies [28, 54], electronic skin [11, 19, 27] and autonomous behavior [35, 40, 52] are enabling robots to engage in increasingly rich, physical interactions with humans: for example, to communicate emotions [16, 78], or to collaborate with humans, such as in object handovers [21, 48, 68]), and caregiving tasks [10, 23]. Conversely, humans may touch robots to express emotions [2], provide guidance [9, 18], or issue instructions. A substantial body of work has investigated affective touch, focusing on both improving the tactile quality of robot touch (e.g., force or temperature [83]) and understanding its psychological and behavioral effects on humans [14, 16, 82]. For instance, studies have shown that robot-initiated touch can foster comfort, increase trust, reduce stress, and enhance perceived social attributes, as well as influence the bonding with the robot [14, 64, 78].

These findings demonstrate the broad potential of physical human-robot interaction in real-world scenarios. But, as robots become increasingly autonomous, driven by advances in AI, human-robot interactions become more dynamic. This raises the need for robots to communicate future touch intent to make touch actions interpretable, efficient and safe. However, beyond verbal cues, which may not always be applicable [43] or effective [6, 10, 67], there is currently a gap in research on how robots communicate touch intent. Parallel work has identified video-based visual feedback as a promising modality for conveying touch intent [81]. We advance this direction through *Mirror Skin*, which leverages the robot's skin as a mirror-like display to communicate *who* shall initiate touch, *where* it will occur, and *when* it is imminent.

2.2 Non-Verbal Robot Intent Communication

To convey robot intent, prior research has explored various forms of non-verbal feedback [51]:

Haptic feedback has been applied to enhance the communication with tele-operated or mobile robots. However, it is typically

limited to conveying predefined states or actions [8, 49], making the communication of richer semantic information challenging.

A very versatile method for conveying robot intent is feedback that the user can visually observe [43]. Research extensively focused on motion-based cues such as eye, head, and arm movements [42, 43, 63]. In particular, robot gaze has been used to increase the clarity of robot actions [43], to improve spatial referencing of objects [42] and humans [41], and to enhance performance in interactive tasks [4, 48]. Moreover, motion trajectories have been optimized for legibility [3, 13], enabling humans to rapidly infer the robot's intended target. Finally, robot gestures can enhance robot communication [6] and prior work explored how to design [32, 63] or generate [44] them.

Other research investigated techniques for visual pre-motion intent communication that allow humans to anticipate the robot's actions early, even before the robot starts moving. Here, light-based approaches have been explored to communicate robot state, motion direction, or intended actions on various types of robots [43, 66, 71, 73]. Nevertheless, light cues have limited semantic capacity, and tasks such as conveying touch locations would require complex mappings, diminishing their effectiveness. Another visual technique is projecting information onto the environment, which has been employed to communicate motion paths [29, 77], state [47], intentions [1] and enhance spatial referencing of objects [25]. However, projection-based methods are constrained by lighting conditions and occlusion due to objects and other entities in the environment [43]. To mitigate these limitations, researchers have utilized augmented reality (AR). For instance to visualize motion direction [24, 76] or trajectories [74, 76]. Moreover, AR has been used to facilitate object handovers [50] and to display robot decision making [49]. But a key limitation of AR is its reliance on additional hardware, restricting its applicability beyond dedicated settings.

Therefore, to convey richer semantic information without additional hardware requirements, one can provide feedback directly on the robot's body. Scholz et al. employed a flexible display wrapped around the robot's arm, demonstrating the feasibility of high-fidelity, body-localized visual feedback [60]. Extending this idea, we envision the robot's skin as a high-resolution display for visual communication, moving beyond simple text or symbolic cues. Inspired by Krüger et al., who used virtual eyes as a mirror to reflect objects [42] or people [41] in order to improve spatial target identification, we leverage the robot's skin as a dynamic mirror that selectively focuses on human body parts to communicate human-robot touch.

2.3 Evaluating Human-Robot Interaction in VR

VR and AR simulations are increasingly used in HRI research due to their effectiveness in evaluating novel robotic system designs through rapid prototyping [45, 80]. Furthermore, prior work has shown that VR-based and physical robot prototypes yield comparable results for pre-touch aspects like visual feedback [55, 62].

Consequently, in this paper, we opt for a two-fold approach: First, to overcome hardware and scalability constraints, we conduct a broad exploration of the *Mirror Skin* concept within a VR environment, alongside a deep investigation of its effectiveness in conveying robot pre-touch intent. Second, we evaluate the UX of *Mirror Skin* during physical interactions using a proof-of-concept prototype in the real world.

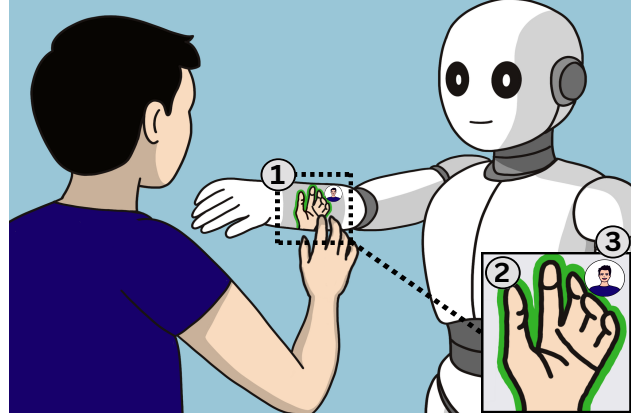


Figure 2: *Mirror Skin*, conveys the contact locations of (1) the robot; (2) the human and (3) who shall initiate touch.

3 Mirror Skin

In the following, we first explain the general concept, as well as a fundamental first version of *Mirror Skin*. Afterwards, we present an iterative design exploration to improve upon the initial version.

3.1 Concept

Similar to cephalopods, we utilize the robot's skin as a high-resolution display conveying visual information. In order to leverage this display for touch intent communication, we must visually convey the three key aspects of the touch interaction: (1) the human body part involved in the touch action, (2) the robot body part involved in the touch action and (3) who shall initiate the touch action.

To encode the location of the intended touch actions, we transform the robot's skin at the location that is involved in the touch interaction into a visual "mirror". This mirror dynamically displays the targeted human body part and its surrounding environment. The mirror reflection continuously follows the motion of the targeted body region, thereby establishing a spatially and temporally aligned one-to-one mapping between the human and robot body. This mirror metaphor leverages humans' familiarity with reflective feedback, which enhances self-identification, and provides an intuitive reference for body-centered interactions [5, 41, 70]. As a result, it enables humans to infer both the occurrence and precise location of the anticipated physical interaction (cf., Figure 2). We selected the body-localized approach over a centralized display (e.g., Pepper's tablet [39]) because this directly encodes the robot's touching body part without requiring an additional mapping. At the same time, it reduces the need to divide attention between the robot's touching body part and the provided visual feedback.

To communicate who shall initiate touch pre-motion, we augment the reflection with actor-specific portraits: A human portrait denotes that the human is expected to touch the robot, whereas a robot portrait indicates that the robot intends to initiate contact.

3.2 Structured Design Exploration

Given that the design space for skin-based communication of robots remains largely unexplored, we performed a structured design exploration in VR to improve the fundamental version of *Mirror Skin* regarding clarity [57] and salience [59]. To identify the key



Figure 3: We investigated six dimensions for visually encoding touch intent on robot skin. For each dimension, we explored multiple design variants to enhance *Mirror Skin*’s clarity and salience and utilize the top-ranked variants for the final design.

parameters that shape the design of *Mirror Skin* we first subdivided the visual feedback of *Mirror Skin* into six dimensions (cf., Figure 3). We then iteratively prototyped within each dimension, generating improved variants informed by prior work and author discussions.

First, to improve the salience and clarity of the content shown in the mirror reflection, we structured the information of the mirror into three dimensions (cf., Figure 3 → Mirror Image Design):

Target encoding specifies how the human body part involved in the touch interaction is represented on the robot’s skin. To provide an unambiguous encoding while minimizing cognitive load, we created several design variants that differ in their degree of abstraction:

DIRECT MIRROR: Our initial version of *Mirror Skin* employs a direct mirror image of the human body part.

OUTLINE & SILHOUETTE: To enhance salience and granularity within the live image, we follow insights from target identification research [53, 75] and apply visual highlighting in the form of an outline or silhouette to the target in the direct mirror.

TEXTURE & COLOR: Inspired by cephalopod camouflage [33, 65] and color communication in HRI [66], we sample the surface texture and color of the target and map it onto the robot’s skin.

SYMBOL: Symbols are an efficient communication modality employed in existing robotic systems [60].

Environmental encoding involves the visual representation of the background in the mirrored image. Although environmental cues can facilitate spatial disambiguation by adding context, visual clutter can negatively impact target identification. Thus, we investigated the following encodings, which gradually reduce the provided environmental information:

FULL ENVIRONMENT: Our initial version of *Mirror Skin* mirrors the full environment onto the robot’s skin.

GRAYSCALE: The environment is desaturated to lower its visual salience while maintaining spatial information.

VIGNETTE: The environment fades out gradually toward the periphery to guide attention to the central target region.

BLUR: The environment is blurred to reduce visual detail while preserving overall scene structure.

NO ENVIRONMENT: The environment is entirely removed to focus solely on the human body.

Spatial encoding specifies how the mirror-image changes based on relative positioning and movement between the human and the robot. By adding spatial cues, we simulate mirror-like behavior, thereby increasing the human’s spatial awareness. We explored the following spatial encoding strategies:

DEPTH: The reflection can either maintain a constant size regardless of distance (no depth) or vary in size relative to the

distance between the human body part and the robot’s surface (with depth), analogous to a real mirror.

DYNAMIC OFFSET: Instead of remaining fixed at the mirror’s center, the reflection is temporarily displaced along the mirror plane in response to human movement and gradually re-centers. This creates the impression that the reflection follows the human, approximating the behavior of a physical mirror.

Second, to convey who shall initiate touch through the mirror, we investigate alternative representations for encoding the initiating actor prior to any robot motion (cf., Figure 3 → Who Touches):

Actor encoding differentiates between the *intention* to touch the human and the *invitation* to be touched by the human. Conveying this distinction clearly can reduce ambiguity in role allocation. We investigated the following encoding strategies:

PORTRAIT: We explored portraits of varying fidelity representing either the human or the robot to indicate the initiator.

SURFACE DEFORMATION: Inspired by prior research [34, 58], we visualize deformations on the robot’s surface similar to how skin gets deformed when touched (i.e., indentations at the front to afford poking, or at the side to afford grasping).

COMPLEMENTARY SHAPE: By altering the shape of the projected area to match the human body part, e.g., a hand-shaped mirror, we create the visual affordance to align the human’s hand with the displayed geometry, conveying an invitation.

CAMERA MOTION: Controlled zoom-in or zoom-out effects to signal whether the robot intends to approach the human or expects the human to initiate contact.

Finally, we sought to explore how to visually integrate *Mirror Skin* on the robot’s skin to make the feedback more noticeable. Thus, we explored (cf., Figure 3 → Skin Integration):

Border encoding determines the visual transition between the reflection and the robot’s body. We explored the following variants:

THIN: The initial version of *Mirror Skin* simply overlays the mirror image onto the skin of the robot, creating a thin transition between the mirror and the remaining robot skin.

THICK: Adding a thick border around the reflection creates a distinct visual boundary, increasing noticeability but reducing the sense of integration with the robot’s body.

SEAMLESS: Seamless embedding of *Mirror Skin* into the robot’s body surface can strengthen the association between the visual feedback and the robot itself. However, increased visual integration may reduce the salience of the reflection.

Attention encoding aims to catch and guide the attention of the human towards the feedback of *Mirror Skin*. This can be particularly helpful when feedback is presented on small robot parts (e.g., fingers), where visual cues may be easily overlooked. To assess the role of attentional guidance, we compare:

NO GUIDANCE: The initial version of *Mirror Skin* relied solely on the salience of the mirror itself to catch the user’s attention.

VISUAL GUIDANCE: The robot provides a salient color cue at a visually prominent region of its body (e.g., the chest) to attract attention. The cue then propagates towards the robot part with *Mirror Skin* feedback, guiding the human’s focus to the mirror.

4 Exploratory Design Study

To evaluate the design variants regarding clarity and salience and to identify further opportunities for improvement, we conducted an exploratory study with robotics and design experts. These insights informed the refined design of *Mirror Skin*.

4.1 Method

Apparatus. To rapidly iterate on the visual encoding strategies, we developed an interactive VR design environment coupled with a GUI. The virtual environment was created with Unity 3D for the Meta Quest 3. We included both a male and female avatar to increase the participants’ self-identification with their avatar, which is critical for attributing the mirrored feedback to one’s body [17]. To meet the quality standards for VR research, we followed previous experiments and used fully rigged avatars from the Microsoft RocketBox library [22] and FinalIK for realistic character animations.

Experimental protocol. Following an initial calibration procedure to align the avatar with the user’s proportions, participants were positioned in front of a humanoid robot (cf., [20, 45]) with integrated *Mirror Skin*. Our setup supported real-time manipulation of all six encoding dimensions, allowing users to select and combine them interactively and observe their effects on the robot’s skin. Furthermore, the environment enabled specification of discrete touch locations and provided predefined poses and motion sequences for both intention and invitation cues, allowing them to examine *Mirror Skin* across varying surface geometries and interaction scenarios.

We first introduced the participants to the different encoding dimensions and our proposed designs. Afterwards, they explored and configured their preferred combination of encoding instantiations and discussed their design choice. We recorded all selections in VR; during the whole experiment, the participants were asked to think aloud while the experimenter took notes. We concluded with a semi-structured interview to discuss participants’ final configurations and to investigate further opportunities for improvement. The study took approximately 60 minutes and was audio-recorded.

Participants. We recruited 7 participants (aged 28 to 34, $\bar{x}$ = 30.71; 5 male, 2 female) with normal or corrected-to-normal vision. Two of them were roboticists with expertise (>5 years) in HRI and human-robot collaboration; two were professional designers (>9 years) with backgrounds in user interface, UX, and interaction design; and the remaining three participants had interdisciplinary expertise (>5 years), combining knowledge from HRI and HCI.

4.2 Results

In the following, we present the expert opinions on our designs and how we refined *Mirror Skin* accordingly.

Target encoding. Participants emphasized the importance of clearly conveying the targeted human body part. In that context, the **DIRECT MIRROR** was not considered clear enough, because “it’s very ambiguous [...] what [the robot] is actually looking at” (P6). Consequently, all participants agreed that the target encoding should be made more salient and identified highlighting (i.e., **OUTLINE** or **SILHOUETTE**) as the most effective approach for the target representation. Moreover, they all favored the **OUTLINE**, because

it retains more detail of the target than the **SILHOUETTE**. While **SYMBOLS** were generally well-received, they lack the dynamic interaction that is available through the mirror: *“It is very static and there is no kind of interaction”* (P4). Additionally, P1 mentioned that *“it’s also a bit more personal to have my moving body displayed here”*, which facilitates the identification between the projected and real body part. Lastly, all participants disliked the plain **TEXTURE & COLOR**, as they are too abstract and ambiguous.

As a result, we chose a live mirror image with an additional **OUTLINE** as our target encoding for *Mirror Skin*.

Environmental encoding. All experts agreed that it is important to reduce the provided background details to avoid visual clutter and they distinguished between two use cases: If it is *“important to interact with objects and environment [...] having these slight environmental cues still helps”* (P1) and therefore, participants chose the **BLUR** feature, as *“blur is suited for blending information without removing it”* (P2). Otherwise, if the robot is solely focused on the human, then it is optimal to remove the background completely (i.e., **NO ENVIRONMENT**), because *“then there’s much clearer focus [...] that it’s about my body”* (P7).

Since the present implementation of *Mirror Skin* is designed to facilitate human-robot touch in single-user scenarios without involving environmental objects, we utilize **NO ENVIRONMENT**, but emphasize the usefulness of **BLUR** for multi-target scenarios.

Spatial encoding. All participants agreed that including **DEPTH** is a good idea, as it resembles the physical property of mirrors and thus helps establish the link between the reflection and our body. The **DYNAMIC OFFSET** feature was mostly preferred (P2, P3, P5-P7), because it contributes to the physical mirror feeling and similarly, P2 stated: *“That looks more organic [...] feels less like tracking”*. Nevertheless, participants (P1, P5, P7) suggested removing the spatial information for body parts with a small form factor like the fingertip, to increase the visibility during motion.

Accordingly, the final version of *Mirror Skin* incorporates both **DEPTH** and **DYNAMIC OFFSET** motion cues for interactions involving larger surface areas, whereas these cues are omitted for small body parts, such as the finger.

Actor encoding. Although approaches like **SURFACE DEFORMATION** and **CAMERA MOTION** were described as *“interesting”* (P6), they were not regarded as intuitive, salient, or efficient to communicate who is going to act. Generally, experts liked the **COMPLEMENTARY SHAPE** (P1, P3, P6-P7), which was considered the most intuitive, naturally inviting humans to initiate touch. But **PORTRAITS** (P1-P7) were clearly regarded as the most efficient. To enhance the interpretability of the portrait, experts (P1, P2, P6) emphasized that it is best to use real headshots of the robot and human to establish a direct connection between the portrait and the actors, similar to a second mirror. That way, the feedback can also convey the actor more easily in multi-user scenarios.

Due to these benefits, we implemented the headshot version of the **PORTRAIT** for *Mirror Skin* and suggest **COMPLEMENTARY SHAPES** as a tool to enhance interpretability during first-time interactions with the robot.

Border encoding. Experts expressed that the border encoding is less critical to its functional design compared to other encoding strategies, as it primarily influences visual aesthetics without substantially enhancing perceptual clarity or salience.

Therefore, we retained the original **THIN** border, which maximizes available screen space relative to the other design alternatives.

Attention encoding. All experts supported the inclusion of an attention-grabbing mechanism (i.e., **VISUAL GUIDANCE**). P4 mentioned: *“I think it’s cool, especially for the small fingertip because otherwise you wouldn’t realize that something is happening there”*. Participants also provided suggestions for improving the implemented version. These included changing the color to be less discouraging for making contact (e.g., blue [66]) (P4) and increasing its pace (P3, P5) to speed up the interaction.

Consequently, we incorporated a fast, uniform blue attention cue for **VISUAL GUIDANCE**.

5 Performance Study

To empirically validate the performance of *Mirror Skin* for communicating robot touch intent across various parts of the upper limbs, we built a VR environment and performed a controlled user study.

Experimental design and task. For our study, we used the VR environment described in Section 4. Participants faced a robot that communicated imminent touch actions, and their task was to infer as quickly as possible *who* should initiate touch and *where* the touch should occur. After pressing a button on the VR controller, which stopped the timer and suppressed all robot feedback to prevent premature responses, participants answered verbally. We restricted interactions to hand-initiated touch, given its predominance in human-robot touch and randomly chose the left or right hand with equal frequency. Likewise, the targeted body part—palm, forearm, or upper arm on either side—was randomly assigned on each trial. We compared three visual communication strategies (cf., Appendix A): First, we investigated *Mirror Skin* for **pre-motion touch intent** communication, meaning that the robot remained in its default pose and communicated intent only through *Mirror Skin*.

Second, we investigated *Mirror Skin* for **in-motion touch intent** communication to assess its performance during dynamic interactions. The robot either offered a body part for contact or approached the participant with the hand to signal touch intent. All approach trajectories followed a direct path to the participant’s body part, with motion and hand orientation dynamically adapting to the participant’s position. The robot always stopped prior to physical contact, ensuring that interactions remained strictly pre-touch.

Finally, to compare *Mirror Skin* to an established visual communication method (cf., [43, 51]), we implemented a baseline condition consisting of the robot motion **gestures** from the previous condition, augmented with **robot gaze** that continuously tracked the relevant human body part to convey explicitly which body part should be involved in the interaction.

Experimental variables. Our study follows a within-subjects design in which each participant interpreted the robots’ communicated touch actions across different types of feedback and action cues. Therefore, we consider two independent variables (IVs):

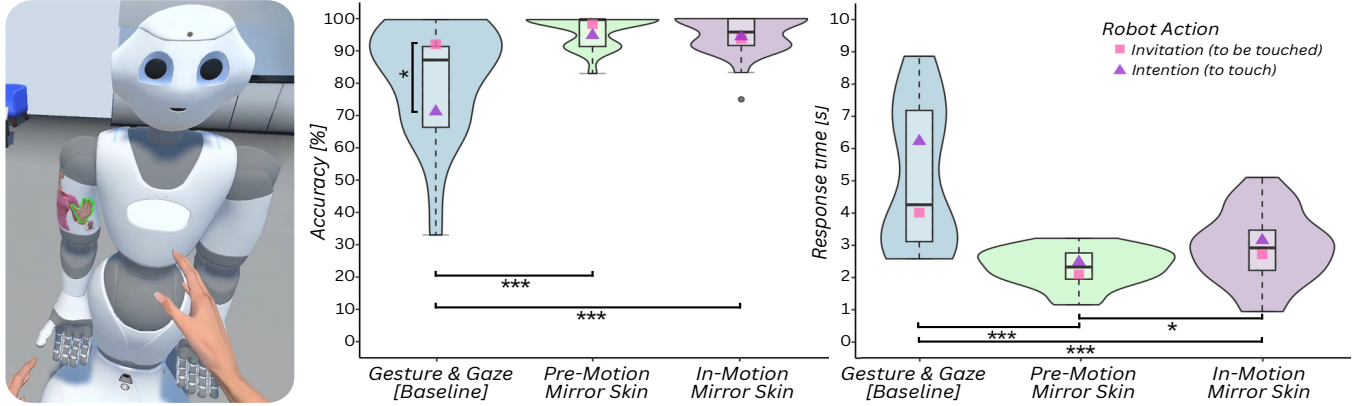


Figure 4: We investigated *Mirror Skin*’s performance for communicating robot touch intent (Left). Our results show that *Mirror Skin* cues before (*pre-motion*) and during (*in-motion*) robot touch motions both significantly enhance accuracy (Middle) and lower response times (Right) of touch intent interpretation, compared to the gesture & gaze baseline.

- **FEEDBACK:** The robot communicates imminent touch actions using either gesture & gaze [BASELINE], pre-motion *Mirror Skin* [PRE-M. MS], or in-motion *Mirror Skin* [IN-M. MS] feedback.
- **ACTION:** The robot can either express the INTENTION to touch the human or offer an INVITATION to be touched.

Each participant completed 12 trials for each ACTION in randomized order, yielding 24 samples per FEEDBACK condition. We counterbalanced the conditions for FEEDBACK with a balanced Latin square, resulting in $3 \times 24 = 72$ samples per participant. For each sample, we measured the following dependent variables (DVs):

- **ACCURACY:** The percentage of correctly answered trials.
- **RESPONSE TIME:** How fast the human identified the robot’s intent.

Procedure. After filling out a consent and demographics form, participants went through the same calibration step as in the exploratory study. To ensure that participants could accurately associate their mirrored body parts with their avatar representation, we added a familiarization phase (cf., [17]), in which participants performed announced and free movements in front of a mirror.

During the study, each FEEDBACK condition was introduced through a conceptual explanation followed by open-ended practice rounds. During the trials, our system tracked participants’ response times and the experimenter documented their answers. After each level of FEEDBACK, participants completed the NASA-TLX [26] to capture their task load, and rated how sufficient and efficient the feedback conveyed the touch intent on 7-point Likert scales.

We concluded with a semi-structured interview for more in-depth insights regarding participants’ opinions and expectations. The study session took approximately 60 minutes and the final interview was audio-recorded.

Participants. We recruited 12 participants (aged 22 to 31, $\bar{x} = 25.16$; 7 male, 5 female) with normal or corrected-to-normal vision and little (4/12) or no experience with HRI (8/12).

Data analysis. We tested the normality assumption of our data with Shapiro-Wilk tests and QQ-plots. Since normality was violated, we applied non-parametric tests to analyze our data: First,

we applied the Aligned Rank Transformation (ART) repeated measures ANOVA as proposed by Wobbrock et al. [79] to investigate interaction and main effects between FEEDBACK and ACTION on accuracy and response time. For significant results, we conducted post hoc analyses using the ART-C procedure, as suggested by Elkin et al. [15]. Effect sizes for ART were reported using partial eta-squared (η_p^2), classified as small ($> .01$), medium ($> .06$), or large ($> .14$) [12]. For ART-C, we reported Cohen’s d , and classified it as small ($> .20$), medium ($> .50$), or large ($> .80$) [12]. Additionally, to analyze the effect of FEEDBACK on task load and UX, we applied Friedman tests, followed by pairwise Wilcoxon signed-rank tests for significant results. We reported Kendall’s W and the rank-biserial r as the measures of effect size and used Cohen’s suggestions [12] to classify them as small ($> .10$), medium ($> .30$), or large ($> .50$).

5.1 Results

Accuracy. Overall, the participants had very high accuracy rates (cf., Figure 4 (middle)), but scores varied between the **BASELINE** (81.93%), **PRE-M. MS** (96.86%) and **IN-M. MS** (94.10%), showing that *Mirror Skin* supported interpreting the robot’s intent. For deeper investigation, we performed an ART that revealed a significant main effect of FEEDBACK on accuracy ($F_{2,55} = 18.38$, $p < .001$, $\eta_p^2 = 0.40$). Post-hoc ART-C pairwise comparisons showed a significantly higher accuracy score for PRE-M. MS compared to BASELINE ($p < .001$, $d = 1.654$) and IN-M. MS also yielded significantly higher accuracy scores than BASELINE ($p < .001$, $d = 1.323$). However, we found no significant difference between the accuracy scores of PRE-M. MS and IN-M. MS ($p > .05$). Next, the ART revealed a significant main effect of ACTION on accuracy, with INVITATION having a significantly higher accuracy compared to INTENTION ($F_{1,55} = 37.52$, $p < .001$, $\eta_p^2 = 0.41$). Additionally, we also found a significant interaction effect between FEEDBACK and ACTION ($F_{2,55} = 11.27$, $p < .001$, $\eta_p^2 = 0.29$). Post-hoc comparisons revealed a significantly higher accuracy for INVITATION compared to INTENTION for the BASELINE ($p < .05$, $d = 1.392$), but showed no significant difference for PRE-M. MS and IN-M. MS ($p > .05$). This observation is supported by post-hoc tests for INTENTION, which

revealed significantly higher accuracy for PRE-M. MS compared to BASELINE ($p < .001$, $d = 1.881$) and for IN-M. MS compared to BASELINE ($p < .001$, $d = 1.933$). Finally, post-hoc tests for INVITATION indicated a significantly higher accuracy for PRE-M. MS compared to BASELINE ($p < .05$, $d = 1.217$), but not between the remaining groups ($p > .05$).

Response time. All types of FEEDBACK conveyed the intent within a reasonable time (cf., Figure 4 (right)); however, with different average response times between **BASELINE (5.12s)**, **PRE-M. MS (2.29s)** and **IN-M. MS (2.93s)**. The ART revealed a significant main effect of FEEDBACK on response time ($F_{2,52} = 34.93$, $p < .001$, $\eta_p^2 = 0.57$). Post-hoc comparisons showed significantly faster response times of PRE-M. MS compared to BASELINE ($p < .001$, $d = -2.433$) and IN-M. MS also yielded significantly faster response times than BASELINE ($p < .001$, $d = -1.642$). Additionally, we found significantly faster response times of PRE-M. MS compared to IN-M. MS ($p < .05$, $d = -0.791$). Next, we found a significant main effect of ACTION on response time, with INVITATION having a significantly faster response time compared to INTENTION ($F_{1,52} = 28.17$, $p < .001$, $\eta_p^2 = 0.35$). Moreover, we also found a significant interaction effect between FEEDBACK and ACTION ($F_{2,52} = 8.25$, $p < .001$, $\eta_p^2 = 0.24$). Contrary to the accuracy, post-hoc tests found no significant differences in response time for INVITATION compared to INTENTION in any Feedback condition ($p > .05$). Furthermore, post-hoc comparisons for INTENTION revealed significantly faster response times for PRE-M. MS compared to BASELINE ($p < .001$, $d = -2.911$) and IN-M. MS also showed significantly faster response times than BASELINE ($p < .01$, $d = -1.747$). Lastly, post-hoc tests for INVITATION indicated significantly faster response times for PRE-M. MS compared to BASELINE ($p < .001$, $d = -2.662$) and for IN-M. MS compared to BASELINE ($p < .01$, $d = -1.641$). These results demonstrate that *Mirror Skin* facilitates faster human comprehension of the robot's intended actions in both pre-motion and in-motion scenarios.

Task load. The task load was similar across BASELINE ($\bar{x} = 37.08$), PRE-M. MS ($\bar{x} = 35.76$) and IN-M. MS ($\bar{x} = 38.47$) and a Friedman test did not find a significant effect of FEEDBACK ($p > .05$).

Usability & User Opinions. First, participants reported that the feedback was sufficient in every condition (BASELINE: $\bar{x} = 6$; PRE-M. MS: $\bar{x} = 6.5$; IN-M. MS: $\bar{x} = 7$) and a Friedman test showed no significant differences ($p > .05$), indicating that the task was comparably feasible with all FEEDBACK types. In line with the prior analysis of response time, participants rated PRE-M. MS ($\bar{x} = 6$) and IN-M. MS ($\bar{x} = 6$) as more efficient for intent communication than BASELINE ($\bar{x} = 4.5$). A Friedman test revealed a significant effect of FEEDBACK on the perceived efficiency ($\chi^2(2) = 8.33$, $p < .05$, $W = 0.347$). Post-hoc tests showed that participants regarded PRE-M. MS as significantly more efficient than BASELINE ($p < .05$, $r = 0.79$), but they found no significant differences between the remaining groups ($p > .05$).

In the follow-up interviews, participants indicated multiple reasons for the improved efficiency. First, several participants (P5, P6, P8, P9, P12) attributed this to the fact that “[one] didn’t have to wait for [the robot’s] movement and gaze” (P9), as all necessary information was immediately available.

At the same time, we observed participants independently leveraging the spatial-referencing advantages of *Mirror Skin* by moving their wrists or fingers to make the mirror feedback more effective (P1, P3-P10, P12). This is notable, as it suggests that the mirroring property offers advantages over static imagery in visual feedback.

Furthermore, all participants (P1-P12) mentioned that when the robot issued an INTENTION, using BASELINE feedback, it was difficult to distinguish between adjacent body parts while the robot’s hand was still further away. In contrast, this issue did not occur in the PRE-M. MS or IN-M. MS conditions, making *Mirror Skin* particularly relevant for robot-initiated touch.

But despite the benefits of *Mirror Skin*, multiple participants (P3-P7, P12) stated that it was more challenging to follow the feedback of *Mirror Skin* during arm motions. This difficulty was repeatedly attributed to self-occlusion (P3, P5, P12), i.e., the robot blocks the line of sight towards the mirror with its own body, and to distortion (P6, P7, P12) caused by motion across curved surfaces, which impaired the clarity and interpretability of the mirrored feedback.

6 Physical Prototype and User Experience Study

To demonstrate the feasibility of *Mirror Skin* on a real robot and to investigate its effects on the UX, we implemented a proof-of-concept *Mirror Skin* prototype and conducted a user study.

6.1 Proof-Of-Concept Prototype

Hardware. We implemented *Mirror Skin* on a humanoid arm-torso setup consisting of two Franka Research FR3 arms from Franka Robotics with seven joints and an Allegro V5 Plus hand from Wonik Robotics with 16 finger joints attached to each arm (cf., Figure 5 (left)). For the experiments, only the right arm was used. To capture the human body parts, we used cameras with a wide field of view (FOV) from InnoMaker (1080p; 130°) on the hand and ELP (2448p; 105°) on the arm. To render the image, we used thin, high-resolution displays from Wisecoco. On the hand, we installed a thin 4-inch LCD (720x720 $\approx$ 255 PPI) and for the arm, we wrapped a flexible 8-inch AMOLED (2460x1860 $\approx$ 385 PPI) around the Franka arm. We used custom 3D-printed mounts to attach the displays and cameras to the corresponding robot part.

Software Pipeline. We implemented the visual mirroring effect in Python, using OpenCV for image capture and processing and Google MediaPipe Landmarks for hand and face tracking. The system includes an initialization procedure that automatically assigns a human actor by detecting their face, capturing a headshot, and storing it as the actor icon.

For each imminent touch interaction, the system allows specification of (1) the robot body part on which the mirror feedback appears, (2) the human body part to be targeted (i.e., left or right hand), and (3) whether the event reflects the robot’s intention to touch or an invitation to be touched. Depending on the selected action, the system displays either the robot or human actor icon.

During execution, the system captures the raw camera frame and detects the targeted human body part. Using the MediaPipe Landmarks, the system then creates an outline around the body part by computing a hand-region mask and drawing its contour. Afterward, it determines the image position of the targeted body part $p_{\{target\}}$, and computes the translation vector $\vec{v}_{\{transl\}}$ required

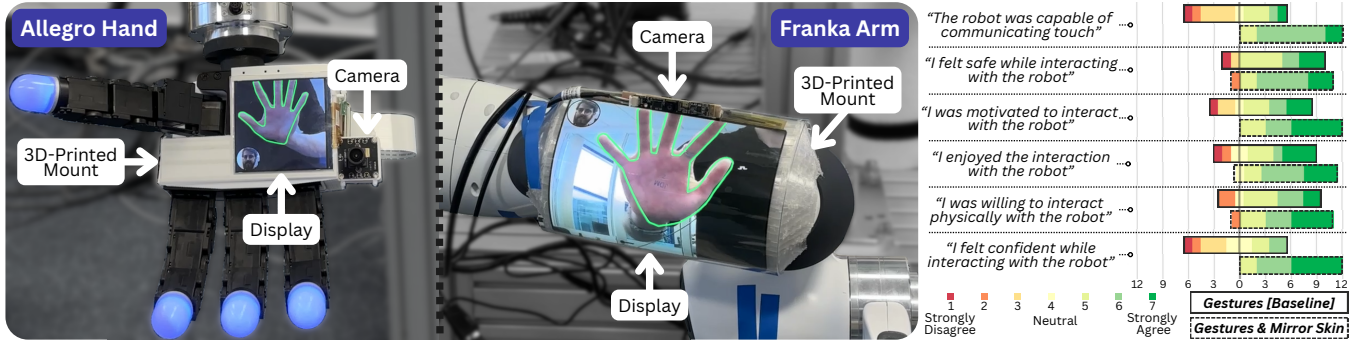


Figure 5: We implemented a proof-of-concept prototype of *Mirror Skin* on a real robot’s forearm and hand (Left) and a user study indicates that our system enhances the user experience compared to only robot gestures (Right).

to center the target in the mirror. To create the dynamic offset, the target is step-wise translated toward the image center at speed s :

$$p_{\{target\}} = p_{\{target\}} + s \cdot \frac{\vec{v}_{\{transl\}}}{|\vec{v}_{\{transl\}}|}$$

Because translating the image introduces empty regions near the camera’s FOV boundaries, we mitigate edge artifacts by cropping a smaller region of the frame and upscaling it to the original resolution through linear interpolation. This creates the illusion of a larger mirror despite the limited FOV. The final frame is then adjusted to the aspect ratio and resolution of the display and rendered as output.

6.2 Method

Experimental design and task. We implemented three physical interaction scenarios: (1) the robot inviting the human to lean onto the robot’s arm for physical support, (2) the robot inviting the participant to place their hand onto the robot’s hand for pulse measurement, and (3) the robot applying a refreshing wipe to the participant’s hand. Prior to each scenario, participants were informed only about the high-level interaction goal (e.g., “the robot aims to measure your pulse”) but were not told what the physical interaction would look like. Their task was to perform this interaction with the robot based on its **FEEDBACK**, consisting of *gestures* [BASELINE], or *gestures & Mirror Skin* [MIRROR SKIN]. The study used a within-subjects design, i.e., participants experienced all scenarios with each feedback type. In addition, the robot targeted either the participant’s left or right hand, and its gestures were adapted accordingly to communicate the intended side. Scenario order was counterbalanced using a balanced Latin square, while the **FEEDBACK** and the target hand were randomized with equal frequency. To ensure safety, the robot operated in a restricted area and both the experimenter and participant had access to an emergency stop.

Procedure. Participants first received a safety briefing and filled out consent and demographic forms. We then explained the *Mirror Skin* concept. Subsequently, participants completed all scenarios, and we recorded whether each interaction was performed successfully, i.e., the participant performed the intended physical interaction with the robot. Afterwards, for each **FEEDBACK** condition, participants filled out the UEQ-S questionnaire and a custom

questionnaire consisting of 7-point Likert-scale items (cf., Figure 5 (right)) adapted from prior HRI studies [61]. We concluded with a semi-structured interview about participants’ preferences and future applications for *Mirror Skin*. Each session took approximately 30 minutes and the interview was audio recorded.

Participants. We recruited 12 participants (aged 23 to 58, $\bar{x} = 37.67$; 9 male, 3 female; 3 left-handed, 9 right-handed). 5 participants had previously experienced some form of physical interaction with a robot, whereas 7 participants had not.

Data analysis. We followed the same steps for data analysis as in study 2 (cf., section 5). For parametric data, we performed paired t -tests and reported Cohen’s d as effect size [12].

6.3 Results

Preferences & User Experience. Although participants successfully completed the tasks in both conditions (BASELINE: 94.4%; MIRROR SKIN: 100.0%), most of them (P1-P2, P4-P11) preferred *Mirror Skin* over the baseline due to its “clear and efficient” (P11) intent communication. Wilcoxon signed-rank tests showed that they rated the robot’s capability to communicate touch intent significantly higher ($p < .05$, $r = 0.76$) for MIRROR SKIN ($\bar{x} = 6$) compared to the BASELINE ($\bar{x} = 4$), as shown in Figure 5 (right). Furthermore, participants reported higher willingness ($p < .05$, $r = 0.70$) to interact physically with the robot for MIRROR SKIN ($\bar{x} = 6$) than for the BASELINE ($\bar{x} = 5$) and they felt more confident during the physical interaction ($p < .01$, $r = 0.86$) when MIRROR SKIN feedback was provided ($\bar{x} = 7$) compared to the BASELINE ($\bar{x} = 4$). Although participants gave MIRROR SKIN high ratings regarding perceived safety ($\bar{x} = 6$), motivation to interact ($\bar{x} = 7$), and enjoyment ($\bar{x} = 6$), these scores did not differ significantly from the BASELINE ($p > .05$).

UEQ-S results show that the pragmatic quality of MIRROR SKIN ($\bar{x} = 1.85$) was significantly higher compared to BASELINE ($\bar{x} = -0.33$), $t(11) = -4.45$, $p < .001$, $d = -1.285$. Also, hedonic quality scores were significantly higher for MIRROR SKIN ($\bar{x} = 1.90$) than for BASELINE ($\bar{x} = 0.19$), $t(11) = -4.17$, $p < .01$, $d = -1.205$. As a result, the overall rating of MIRROR SKIN ($\bar{x} = 1.88$) was significantly higher than BASELINE ($\bar{x} = -0.07$), $t(11) = -4.70$, $p < .001$, $d = -1.359$. Relative to a benchmark dataset of 468 studies [30], these ratings place *Mirror Skin* in the “Excellent” UX category.

Opportunities & Applicability. Participants saw potential for *Mirror Skin* across diverse pHRI domains such as medical and caregiving contexts (P1-P3, P6-P11), collaborative and industrial settings (P1, P3-P5, P8-P9, P11-P12), educational scenarios (P2-P3, P10, P12), and home or service robots (P6, P8-P9, P11).

Especially in collaborative scenarios like object handovers, participants (P1, P4, P5, P8, P11) saw opportunities for rich semantic feedback, as current gestural interactions are not precise enough. P4 noted: “If [the robot] tries to hand over something [...] you put two hands up to see which hand [the robot] wants to touch”. *Mirror Skin*’s rich semantic feedback, which tells exactly where contact happens, would reduce ambiguity in these scenarios.

Furthermore, participants indicated that *Mirror Skin* is particularly suited for initial interactions with robots, as the feedback was found to be “very intuitive” (P1) and “exciting” (P5), which encourages humans to interact physically with the robot. P8 noted that this may be particularly beneficial for robot-assisted medical treatments, for example “in physiotherapy, during the massage, it gives you more security or more confidence, because you have confirmation that the robot aims for the right body part”.

Moreover, P3 hypothesized that *Mirror Skin* is particularly important for robots with non-humanoid morphology, as “it’s easier to communicate with *Mirror Skin*, ‘touch this!’ than [verbally describe] whatever is between the third and the fourth joint”, thus indicating advantages of rich visual feedback over other modalities. Finally, participants (P5, P8, P11, P12) suggested extending the mirror metaphor to signal points of interest, for example, communicating the objects a robot plans to grasp.

7 Discussion & Future Directions

In this section, we interpret our findings and outline implications for advancing visual skin-based communication. We also discuss the limitations of this work and identify directions for future research.

7.1 Enhancing Robot Touch Intent Communication with *Mirror Skin*

Our results demonstrate that *Mirror Skin* enables faster and more accurate interpretation of touch intent compared to gestures and gaze. One key advantage of visual, skin-based communication is its ability to convey all relevant information at a glance, without requiring robot motion. This enables users to anticipate forthcoming actions and thus could be an efficient mechanism for safe pHRI.

We also found that gesture-based communication (e.g., pointing) offers limited spatial precision, particularly when the robot’s end effector is farther away from the human. In contrast, *Mirror Skin* provides fine-grained spatial feedback, which is well-suited for conveying precise touch points. Furthermore, unlike other feedback methods (e.g., gaze) which are typically limited to a single point of focus, *Mirror Skin* supports the simultaneous visualization of multiple touch intentions. Future work should investigate this spatial multiplexing as it offers interesting opportunities for communicating touch intent in multi-contact or multi-user scenarios.

Finally, we showed that clearer articulation of the robot’s intent increases humans’ willingness and confidence to engage in physical contact, highlighting the value of higher-semantic feedback. As robots increasingly perform physical interactions in domains

such as healthcare and industry collaboration, fostering human acceptance of robot touch becomes crucial.

7.2 Expanding *Mirror Skin*’s Semantic Information for Robot Touch

Similar to the visual feedback of cephalopods, *Mirror Skin* is not constrained in what information can be visually mapped onto the robot’s skin. A natural extension of communicating touch actions is to explore how *Mirror Skin* can be leveraged to convey information about the quality of the touch, such as the robot’s precise touch action and the intensity of contact. We hypothesize that providing humans with anticipatory information about how the robot will touch their body may enhance comfort and trust.

Another promising direction is to extend *Mirror Skin* to support bidirectional communication, enabling the robot to respond to human touch intent. For instance, when a human hand approaches the robot’s arm, *Mirror Skin* could be used to signal the robot’s awareness of the incoming interaction and visually indicate acceptance, rejection, or suggest a counterproposal. When a human hand approaches a sensitive or safety-critical region of the robot’s body, *Mirror Skin* could be employed to deny the interaction and redirect it to a safer adjacent location. Given the high spatial accuracy and rapid interpretability of *Mirror Skin*, such responsive feedback could facilitate real-time negotiation of touch interactions, enhancing both safety and fluidity in human-robot collaboration.

7.3 *Mirror Skin* for Diverse Robot Communication and Morphologies

While our work focused specifically on touch intent, the underlying principles do apply to other forms of human-robot communication. For instance, building on the work of Krüger et al. [41, 42], *Mirror Skin* could be extended to support spatial referencing, such as indicating the robot’s intention to grasp a specific object. Crucially, unlike prior mirroring approaches that centralize feedback to one specific region (e.g., the robot’s eyes), *Mirror Skin* localizes the visual feedback directly at the relevant body part. This eliminates the need for humans to divide attention between multiple regions (e.g., gaze and end effector), potentially increasing efficiency (cf., [73]).

Moreover, we see potential in high-resolution, skin-based visual feedback, including textures, images and videos, as a promising expressive medium for robots. Building on the work of Hu et al. [34], such visual feedback could convey messages to the human, indicate internal robot state, or communicate emotional affect during social human-robot interactions.

Furthermore, to simplify the interpretation of gestures and gaze, we restricted touch actions to be performed with the hand. However, in real-world scenarios, robots may initiate contact using other body parts, for example, lifting a person with the forearm or stabilizing with the torso. *Mirror Skin* appears to be well-suited to communicate such interactions, as it is inherently morphology-independent. This also offers a promising communication modality for appearance-constrained robots (e.g., [7, 58, 66, 76]) that lack anthropomorphic features and hence cannot rely on conventional nonverbal cues such as gestures or gaze.

Finally, although we attached all screens for *Mirror Skin*, numerous existing robotic platforms already incorporate body-mounted

displays, particularly on the torso [39] and face [37]. These existing displays could be directly repurposed for *Mirror Skin*, or supplemented with additional screens to extend coverage. As skin-based display technology matures, we envision the entire robot body eventually functioning as a continuous display surface (cf., [38]).

7.4 Limitations & Future Work

While *Mirror Skin* demonstrates strong potential for communicating touch intent, our current implementation is subject to several limitations that we plan to address in future work.

Advancing the *Mirror Skin* prototype. In our current implementation of *Mirror Skin*, self-occlusion remains an open challenge, particularly in dynamic scenarios where the robot’s own body may obstruct the visual feedback. In future work, we intend to address this limitation through several approaches. First, occlusion effects could be mitigated by rendering *Mirror Skin* on adjacent, non-occluded surface regions while maintaining a visual reference to the originally targeted surface area. Moreover, backwards-facing cameras could be integrated to simulate visual transparency of occluding body parts (cf., [72]). Finally, supplementary hardware such as AR headsets could be employed to render visual feedback directly above or through occluding body parts (cf., [24]).

At the same time, our proof-of-concept was limited to only one hand and forearm of the robot. In addition, although we used wide-angle cameras, they did not cover every angle during robot motion. Future work should equip a robot with more displays and cameras to remove blind spots and evaluate *Mirror Skin* on a fully scaled robot.

Tailor *Mirror Skin* to dedicated settings and scenarios. First, the current design of *Mirror Skin* has been evaluated with able-bodied adults. However, key application domains such as caregiving often involve user groups with varying perceptual and cognitive profiles, including individuals with visual impairments. Consequently, future work should conduct a stakeholder analysis to identify the specific needs of these user groups and explore adaptations of the concept to ensure inclusive communication.

Second, to remain within a feasible scope, we evaluated *Mirror Skin* only in a controlled, safe lab environment. Future work should investigate its performance and cognitive load across a broader range of contexts and tasks, and adapt *Mirror Skin*’s visual feedback accordingly. For example, during repetitive or recurrent interactions, the feedback could be simplified and made more static (e.g., by removing depth and dynamic offset) to reduce visual load. Additionally, although we compared *Mirror Skin* with gestures and gaze as established non-verbal communication modalities, future research should explore its integration into multimodal robotic systems to potentially further advance communicative efficiency.

Third, our findings suggest that *Mirror Skin* is particularly well-suited to initial human-robot interactions, as it appears to encourage physical engagement with the robot. A next step will be to investigate the effects of *Mirror Skin* during long-term use, once the potential novelty effect has diminished. To sustain engagement and efficiency over long-term use, it may be advantageous to grant users control over customizing the visual feedback along our proposed design dimensions.

Fourth, *Mirror Skin* requires live camera images of people and the environment, which could raise privacy concerns for robots in public spaces, especially regarding bystanders who are not actively interacting with the robot. Therefore, we emphasize techniques like background removal and face blurring to mitigate these issues and encourage further investigation in future work.

Finally, the potential intuition of the *Mirror Skin* deserves further investigation. The mirror metaphor appears to facilitate the understanding and quick use of the displayed information. We observed that participants actively utilized laws governing the use of actual mirrors for disambiguation, e.g., by moving their wrist or fingers to identify the targeted hand. However, our evaluation of *Mirror Skin* was conducted in scenarios where humans were explicitly introduced to the interface concept and its semantics. Future work should investigate *Mirror Skin*’s use without such guidance and, if required, identify further means for promoting intuitive use.

8 Conclusion

In this work, we introduce *Mirror Skin*, a cephalopod-inspired concept for conveying robot touch intent via high-resolution, mirror-like visual feedback on robot skin. By mapping visual representations of human body parts onto the robot’s touching surface, *Mirror Skin* communicates *who* shall initiate touch, *where* it will occur and *when* it is imminent. We developed and evaluated the *Mirror Skin* concept through a mixed methods approach consisting of a structured design exploration in VR, a proof-of-concept prototype and three controlled user studies. Our results demonstrate that *Mirror Skin* significantly improves both accuracy and response time in interpreting touch intent compared to a gesture-and-gaze baseline and that *Mirror Skin* enhances the UX, including greater willingness and confidence to engage in pHRI. *Mirror Skin* offers new avenues for robot skin-based visual interfaces, dynamic bidirectional human-robot interaction, multi-target touch communication, and morphology-independent feedback. This makes *Mirror Skin* a promising concept for various robotic platforms and applications.

Acknowledgments

We thank all participants of our user studies, as well as Oliver Schön, for his support during the fabrication of the *Mirror Skin* prototype. Moreover, we thank Hannah Wagmann for sketching our conceptual space. We also thank the reviewers for their valuable comments.

References

- [1] Rasmus S. Andersen, Ole Madsen, Thomas B. Moeslund, and Heni Ben Amor. 2016. Projecting robot intentions into human environments. In *2016 25th IEEE International Symposium on Robot and Human Interactive Communication (RO-MAN)*. 294–301. doi:10.1109/ROMAN.2016.7745145
- [2] Rebecca Andreasson, Beatrice Alenljung, Erik Billing, and Robert Lowe. 2018. Affective Touch in Human–Robot Interaction: Conveying Emotion to the Nao Robot. *International Journal of Social Robotics* 10, 4 (Sept. 2018), 473–491. doi:10.1007/s12369-017-0446-3
- [3] Christopher Bodden, Daniel Rakita, Bilge Mutlu, and Michael Gleicher. 2018. A flexible optimization-based method for synthesizing intent-expressive robot arm motion. *The International Journal of Robotics Research* 37, 11 (Sept. 2018), 1376–1394. doi:10.1177/0278364918792295 Publisher: SAGE Publications Ltd STM.
- [4] Jean-David Boucher, Ugo Pattacini, Amelie Lelong, Gerard Bailly, Frederic Elisei, Sascha Fagel, Peter F. Dominey, and Jocelyne Ventre-Dominey. 2012. I Reach Faster When I See You Look: Gaze Effects in Human–Human and Human–Robot

- Face-to-Face Cooperation. *Frontiers in Neurorobotics* Volume 6 - 2012 (2012). doi:10.3389/fnbot.2012.00003
- [5] Alexandria Boyle. 2018. Mirror Self-Recognition and Self-Identification. *Philosophy and Phenomenological Research* 97, 2 (2018), 284–303.
 - [6] C. Breazeal, C.D. Kidd, A.L. Thomaz, G. Hoffman, and M. Berlin. 2005. Effects of nonverbal communication on efficiency and robustness in human-robot teamwork. In *2005 IEEE/RSJ International Conference on Intelligent Robots and Systems*. 708–713. doi:10.1109/IROS.2005.1545011
 - [7] Jessica R. Cauchard, Charles Dutau, Gianluca Corsini, Marco Cognetti, Daniel Sidobre, Simon Lacroix, and Anke M. Brock. 2024. Considerations for Handover and Co-working with Drones. In *Companion Proceedings of the 2024 ACM/IEEE International Conference on Human-Robot Interaction* (Boulder, CO, USA) (*HRI Companion '24*). Association for Computing Machinery, New York, NY, USA, 302–306. doi:10.1145/3610978.3640590
 - [8] Yuhang Che, Cuthbert T. Sun, and Allison M. Okamura. 2018. Avoiding Human-Robot Collisions Using Haptic Communication. In *2018 IEEE International Conference on Robotics and Automation (ICRA)*. 5828–5834. doi:10.1109/ICRA.2018.8460946 ISSN: 2577-087X.
 - [9] Tiffany L. Chen and Charles C. Kemp. 2010. Lead me by the hand: Evaluation of a direct physical interface for nursing assistant robots. In *2010 5th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*. 367–374. doi:10.1109/HRI.2010.5453162
 - [10] Tiffany L. Chen, Chih-Hung King, Andrea L. Thomaz, and Charles C. Kemp. 2011. Touched by a robot: an investigation of subjective responses to robot-initiated touch. In *Proceedings of the 6th International Conference on Human-Robot Interaction* (Lausanne, Switzerland) (*HRI '11*). Association for Computing Machinery, New York, NY, USA, 457–464. doi:10.1145/1957656.1957818
 - [11] Gordon Cheng, Emmanuel Dean-Leon, Florian Bergner, Julio Rogelio Guadarrama Olvera, Quentin Lebouet, and Philipp Mittendorf. 2019. A Comprehensive Realization of Robot Skin: Sensors, Sensing, Control, and Applications. *Proc. IEEE* 107, 10 (Oct. 2019), 2034–2051. doi:10.1109/JPROC.2019.2933348
 - [12] Jacob Cohen. 1988. *Statistical Power Analysis for the Behavioral Sciences*. Routledge. doi:10.4324/9780203771587
 - [13] Anca D. Dragan, Shira Bauman, Jodi Forlizzi, and Siddhartha S. Srinivasa. 2015. Effects of Robot Motion on Human-Robot Collaboration. In *Proceedings of the Tenth Annual ACM/IEEE International Conference on Human-Robot Interaction* (Portland, Oregon, USA) (*HRI '15*). Association for Computing Machinery, New York, NY, USA, 51–58. doi:10.1145/2696454.2696473
 - [14] Monika Eckstein, Ilshat Mamaev, Beate Ditzgen, and Uta Sailer. 2020. Calming effects of touch in human, animal, and robotic interaction—scientific state-of-the-art and technical advances. *Frontiers in psychiatry* 11 (2020), 555058. doi:10.3389/fpsy.2020.555058
 - [15] Lisa A. Elkin, Matthew Kay, James J. Higgins, and Jacob O. Wobbrock. 2021. An Aligned Rank Transform Procedure for Multifactor Contrast Tests. In *The 34th Annual ACM Symposium on User Interface Software and Technology* (Virtual Event, USA) (*UIST '21*). Association for Computing Machinery, New York, NY, USA, 754–768. doi:10.1145/3472749.3474784
 - [16] Jan B.F. Van Erp and Alexander Toet. 2013. How to Touch Humans: Guidelines for Social Agents and Robots That Can Touch. In *2013 Humaine Association Conference on Affective Computing and Intelligent Interaction*. 780–785. doi:10.1109/ACII.2013.145
 - [17] Martin Feick, André Zenner, Simon Seibert, Anthony Tang, and Antonio Krüger. 2024. The Impact of Avatar Completeness on Embodiment and the Detectability of Hand Redirection in Virtual Reality. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (*CHI '24*). Association for Computing Machinery, New York, NY, USA, Article 548, 9 pages. doi:10.1145/3613904.3641933
 - [18] Kerstin Fischer, Franziska Kirstein, Lars Christian Jensen, Norbert Krüger, Kamil Kukliński, Maria Vanessa aus der Wieschen, and Thiusius Rajeeth Savarimuthu. 2016. A comparison of types of robot control for programming by Demonstration. In *2016 11th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*. 213–220. doi:10.1109/HRI.2016.7451754
 - [19] Xuemei Fu, Wen Cheng, Guanxiang Wan, Zijie Yang, and Benjamin C. K. Tee. 2024. Toward an AI Era: Advances in Electronic Skins. *Chemical Reviews* 124, 17 (Sept. 2024), 9899–9948. doi:10.1021/acs.chemrev.4c00049
 - [20] Amit Ghimire, Anova Hou, Ig-Jae Kim, and Dongwook Yoon. 2025. AvatARoid: A Motion-Mapped AR Overlay to Bridge the Embodiment Gap Between Robots and Teleoperators in Robot-Mediated Telepresence. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems* (*CHI '25*). Association for Computing Machinery, New York, NY, USA, Article 906, 26 pages. doi:10.1145/3706598.3713812
 - [21] Michael Gienger, Dirk Ruiken, Tamas Bates, Mohamed Regaieg, Michael Meibner, Jens Kober, Philipp Seiwald, and Arne-Christoph Hildebrandt. 2018. Human-robot cooperative object manipulation with contact changes. In *2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 1354–1360. doi:10.1109/IROS.2018.8594140
 - [22] Mar Gonzalez-Franco, Eyal Ofek, Ye Pan, Angus Antley, Anthony Steed, Bernhard Spanlang, Antonella Maselli, Domna Banakou, Nuria Pelechano, Sergio Orts-Escobedo, Veronica Orvalho, Laura Trutoiu, Markus Wojcik, Maria V. Sanchez-Vives, Jeremy Bailenson, Mel Slater, and Jaron Lanier. 2020. The Rocketbox Library and the Utility of Freely Available Rigged Avatars. *Frontiers in Virtual Reality* 1 (Nov. 2020). doi:10.3389/frvir.2020.561558 Publisher: Frontiers.
 - [23] Carina Soledad González-González, Verónica Violant-Holz, and Rosa Maria Gil-Iranzo. 2021. Social Robots in Hospitals: A Systematic Review. *Applied Sciences* 11, 13 (2021). doi:10.3390/app11135976
 - [24] Morris Gu, Akansel Cosgun, Wesley P. Chan, Tom Drummond, and Elizabeth Croft. 2021. Seeing Thru Walls: Visualizing Mobile Robots in Augmented Reality. In *2021 30th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*. 406–411. doi:10.1109/RO-MAN50785.2021.9515322
 - [25] Zhao Han, Alexander Wilkinson, Jenna Parrillo, Jordan Allspaw, and Holly A. Yanco. 2020. Projection Mapping Implementation: Enabling Direct Externalization of Perception Results and Action Intent to Improve Robot Explainability. doi:10.48550/arXiv.2010.02263 arXiv:2010.02263 [cs].
 - [26] Sandra G Hart and Lowell E Staveland. 1988. Development of NASA-TLX (Task Load Index): Results of Empirical and Theoretical Research. In *Advances in Psychology*. Vol. 52. Elsevier, 139–183.
 - [27] Yin He, Xiaoxuan Xu, Shuang Xiao, Junxian Wu, Peng Zhou, Li Chen, and Hao Liu. 2024. Research Progress and Application of Multimodal Flexible Sensors for Electronic Skin. *ACS Sensors* 9, 5 (May 2024), 2275–2293. doi:10.1021/acssensors.4c00307
 - [28] Chidanand Hegde, Jiangtao Su, Joel Ming Rui Tan, Ke He, Xiaodong Chen, and Shlomo Magdassi. 2023. Sensing in soft robotics. *ACS nano* 17, 16 (2023), 15277–15307. doi:10.1021/acsnano.3c04089
 - [29] Nicholas J. Hetherington, Elizabeth A. Croft, and H.F. Machiel Van der Loos. 2021. Hey Robot, Which Way Are You Going? Nonverbal Motion Legibility Cues for Human-Robot Spatial Interaction. *IEEE Robotics and Automation Letters* 6, 3 (2021), 5010–5015. doi:10.1109/LRA.2021.3068708
 - [30] Andreas Hinderks, Martin Schrepp, and Jörg Thomaschewski. 2018. A benchmark for the short version of the user experience questionnaire. In *Proceedings of the 14th International Conference on Web Information Systems and Technologies-APMD WE*. SciTePress, 373–377.
 - [31] Takahiro Hirano, Masahiro Shiomi, Takamasa Iio, Mitsuhiro Kimoto, Takuya Nagashio, Ivan Tanev, Katsunori Shimohara, and Norihiro Hagita. 2016. Communication Cues in a Human-Robot Touch Interaction. In *Proceedings of the Fourth International Conference on Human Agent Interaction* (Biopolis, Singapore) (*HAI '16*). Association for Computing Machinery, New York, NY, USA, 201–206. doi:10.1145/2974804.2974809
 - [32] Rachel M. Holladay, Anca D. Dragan, and Siddhartha S. Srinivasa. 2014. Legible robot pointing. In *The 23rd IEEE International Symposium on Robot and Human Interactive Communication*. 217–223. doi:10.1109/ROMAN.2014.6926256
 - [33] Martin J. How, Mark D. Norman, Julian Finn, Wen-Sung Chung, and N. Justin Marshall. 2017. Dynamic Skin Patterns in Cephalopods. *Frontiers in Physiology* 8 (June 2017). doi:10.3389/fphys.2017.00393 Publisher: Frontiers.
 - [34] Yuhang Hu and Guy Hoffman. 2023. What Can a Robot's Skin Be? Designing Texture-changing Skin for Human-Robot Social Interaction. *J. Hum.-Robot Interact.* 12, 2, Article 26 (April 2023), 19 pages. doi:10.1145/3532772
 - [35] Yiyang Huang, Yuhui Hao, Bo Yu, Feng Yan, Yuxin Yang, Feng Min, Yinhe Han, Lin Ma, Shaoshan Liu, Qiang Liu, and Yiming Gan. 2025. Dadu-Corki: Algorithm-Architecture Co-Design for Embodied AI-powered Robotic Manipulation. In *Proceedings of the 52nd Annual International Symposium on Computer Architecture (ISCA '25)*. Association for Computing Machinery, New York, NY, USA, 327–343. doi:10.1145/3695053.3731099
 - [36] Pierce Hutton, Brett M Seymoure, Kevin J McGraw, Russell A Ligon, and Richard K Simpson. 2015. Dynamic color communication. *Current Opinion in Behavioral Sciences* 6 (Dec. 2015), 41–49. doi:10.1016/j.cobeha.2015.08.007
 - [37] Alisa Kalebina, Grace Schroeder, Aidan Allchin, Kearsa Berlin, and Maya Cakmak. 2018. Characterizing the Design Space of Rendered Robot Faces. In *Proceedings of the 2018 ACM/IEEE International Conference on Human-Robot Interaction* (Chicago, IL, USA) (*HRI '18*). Association for Computing Machinery, New York, NY, USA, 96–104. doi:10.1145/3171221.3171286
 - [38] Kiwook Kim, Dong Ryong Kim, Dohyeon Kim, Hyeon Hwa Song, Seungmin Lee, Yonghoon Choi, Kyunghoon Lee, Gwang Heon Lee, Jinhee Lee, Hye Hyun Kim, Eonhyoung Ahn, Jae Hong Jang, Jaewon Kim, Hyo Cheol Lee, Yunho Kim, Soo Ik Park, Jisu Yoo, Youngsik Lee, Jongnam Park, Dae-Hyeon Kim, Moon Kee Choi, and Jiwoong Yang. 2025. Intrinsically-Stretchable and Patternable Quantum Dot Color Conversion Layers for Stretchable Displays in Robotic Skin and Wearable Electronics. *Advanced Materials* 37, 32 (2025), 2420633. doi:10.1002/adma.202420633
 - [39] Yujin Kim, Christine P Lee, and Bilge Mutlu. 2026. Speaking with Screens: Design Space and Guidelines for Informational Robot Screens. In *Proceedings of the 21st ACM/IEEE International Conference on Human-Robot Interaction* (Edinburgh, Scotland, UK) (*HRI '26*). Association for Computing Machinery, New York, NY, USA, 385–394. doi:10.1145/3757279.3785557

- [40] Susanna Krämer, Eyal Ginosar, Marion Koelle, Heiko Müller, Susanne Boll, and Jessica Cauchard. 2025. “I Have a Package For Mr. and Mrs. Smith, Can You Point Me to Their House?”—Co-Designing Bystanders’ Interaction with Autonomous Delivery Vehicles. *J. Hum.-Robot Interact.* 14, 3, Article 53 (May 2025), 29 pages. doi:10.1145/3724131
- [41] Matti Krüger, Yutaka Oshima, and Yu Fang. 2026. Virtual Reflections on a Dynamic 2D Eye Model Improve Spatial Reference Identification. *IEEE Transactions on Human-Machine Systems* (2026), 1–10. doi:10.1109/THMS.2026.3651818
- [42] Matti Krüger, Daniel Tanneberg, Chao Wang, Stephan Hasler, and Michael Gienger. 2025. Mirror Eyes: Explainable Human-Robot Interaction at a Glance. In *2025 34th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*. 797–804. doi:10.1109/RO-MAN63969.2025.11217810
- [43] Gregory Lemasurier, Gal Bejerano, Victoria Albanese, Jenna Parrillo, Holly A. Yanco, Nicholas Amerson, Rebecca Hetrick, and Elizabeth Phillips. 2021. Methods for Expressing Robot Intent for Human-Robot Collaboration in Shared Workspaces. *J. Hum.-Robot Interact.* 10, 4, Article 40 (Sept. 2021), 27 pages. doi:10.1145/3472223
- [44] Jan Leusmann, Steeven Villa, Thomas Liang, Chao Wang, Albrecht Schmidt, and Sven Mayer. 2025. An Approach to Elicit Human-Understandable Robot Expressions to Support Human-Robot Interaction. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI '25)*. Association for Computing Machinery, New York, NY, USA, Article 534, 17 pages. doi:10.1145/3706598.3713085
- [45] Hanfang Lyu, Xiaoyu Wang, Nandi Zhang, Shuai Ma, Qian Zhu, Yuhuan Luo, Fugee Tsung, and Xiaojuan Ma. 2025. Signaling Human Intentions to Service Robots: Understanding the Use of Social Cues during In-Person Conversations. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI '25)*. Association for Computing Machinery, New York, NY, USA, Article 603, 21 pages. doi:10.1145/3706598.3714235
- [46] Eloise Matheson, Riccardo Minto, Emanuele G. G. Zampieri, Maurizio Faccio, and Giulio Rosati. 2019. Human-Robot Collaboration in Manufacturing Applications: A Review. *Robotics* 8, 4 (2019). doi:10.3390/robotics8040100
- [47] Takafumi Matsumaru. 2006. Mobile Robot with Preliminary-announcement and Display Function of Forthcoming Motion using Projection Equipment. In *ROMAN 2006 - The 15th IEEE International Symposium on Robot and Human Interactive Communication*. 443–450. doi:10.1109/ROMAN.2006.314368
- [48] Ajung Moon, Daniel M. Troniak, Brian Gleeson, Matthew K.X.J. Pan, Minhua Zheng, Benjamin A. Blumer, Karon MacLean, and Elizabeth A. Croft. 2014. Meet me where i’m gazing: how shared attention gaze affects human-robot handover timing. In *Proceedings of the 2014 ACM/IEEE International Conference on Human-Robot Interaction (Bielefeld, Germany) (HRI '14)*. Association for Computing Machinery, New York, NY, USA, 334–341. doi:10.1145/2559636.2559656
- [49] James F. Mullen, Josh Mosier, Sounak Chakrabarti, Anqi Chen, Tyler White, and Dylan P. Losey. 2021. Communicating Inferred Goals With Passive Augmented Reality and Active Haptic Feedback. *IEEE Robotics and Automation Letters* 6, 4 (Oct. 2021), 8522–8529. doi:10.1109/LRA.2021.3111055
- [50] Rhys Newbury, Akansel Cosgun, Tysha Crowley-Davis, Wesley P. Chan, Tom Drummond, and Elizabeth A. Croft. 2022. Visualizing Robot Intent for Object Handovers with Augmented Reality. In *2022 31st IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*. 1264–1270. doi:10.1109/RO-MAN53752.2022.9900524
- [51] Max Pascher, Uwe Gruenefeld, Stefan Schneegass, and Jens Gerken. 2023. How to Communicate Robot Motion Intent: A Scoping Review. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems (Hamburg, Germany) (CHI '23)*. Association for Computing Machinery, New York, NY, USA, Article 409, 17 pages. doi:10.1145/3544548.3580857
- [52] Hannah R. M. Pelikan, Stuart Reeves, and Marina N. Cantarutti. 2024. Encountering Autonomous Robots on Public Streets. In *Proceedings of the 2024 ACM/IEEE International Conference on Human-Robot Interaction (Boulder, CO, USA) (HRI '24)*. Association for Computing Machinery, New York, NY, USA, 561–571. doi:10.1145/3610977.3634936
- [53] Loizos Psarakis, Dimitris Nathanael, and Nicolas Marmaras. 2022. Fostering short-term human anticipatory behavior in human-robot collaboration. *International Journal of Industrial Ergonomics* 87 (2022), 103241. doi:10.1016/j.ergon.2021.103241
- [54] Peter Roberts, Mason Zadan, and Carmel Majidi. 2021. Soft tactile sensing skins for robotics. *Current Robotics Reports* 2, 3 (2021), 343–354. doi:10.1007/s43154-021-00065-2
- [55] Juan-Carlos Rojas, Jaime Alvarez, Arantza Garcia-Mora, and Paulina Méndez. 2024. Comparison of Robot Assessment by Using Physical and Virtual Prototypes: Assessment of Appearance Characteristics, Emotional Response and Social Perception. In *International Conference on Human-Computer Interaction*. Springer, 127–145.
- [56] Matteo Rubagotti, Inara Tusseyeva, Sara Baltabayeva, Danna Summers, and Anara Sandygulova. 2022. Perceived safety in physical human-robot interaction—A survey. *Robotics and Autonomous Systems* 151 (2022), 104047. doi:10.1016/j.robot.2022.104047
- [57] Roya Salehzadeh, Jiaqi Gong, and Nader Jalili. 2022. Purposeful Communication in Human-Robot Collaboration: A Review of Modern Approaches in Manufacturing. *IEEE Access* 10 (2022), 129344–129361. doi:10.1109/ACCESS.2022.3227049
- [58] Bob R. Schadenberg, Jorge Davo Sainz, Jan Kolkmeier, Hideki Garcia Goo, and Vanessa Evers. 2025. Goo-y: A Robot’s Shape-Changing Capability for Fast, Non-Anthropomorphic Communication with People (and Animals). In *Proceedings of the Extended Abstracts of the CHI Conference on Human Factors in Computing Systems (CHI EA '25)*. Association for Computing Machinery, New York, NY, USA, Article 308, 6 pages. doi:10.1145/3706599.3719793
- [59] Guido Schillaci, Saša Bodiřoza, and Verena Vanessa Hafner. 2013. Evaluating the effect of saliency detection and attention manipulation in human-robot interaction. *International Journal of Social Robotics* 5, 1 (2013), 139–152. doi:10.1007/s12369-012-0174-7
- [60] Constantin Scholz, Hoang-Long Cao, Ilias El Makrini, Susanne Niehaus, Maximilian Kaufmann, David Cheyns, Nima Roshandel, Aleksander Burkiewicz, Mariane Shhaitly, Emil Imrith, Jan Genoe, Xavier Rottenberg, Peter Gerets, and Bram Vanderborght. 2025. Improving robot-to-human communication using flexible display technology as a robotic-skin-interface: a co-design study. *International Journal of Intelligent Robotics and Applications* 9, 1 (March 2025), 146–163. doi:10.1007/s43155-024-00343-0
- [61] Hasti Seifi, Arpit Bhatia, and Kasper Hornbæk. 2024. Charting User Experience in Physical Human-Robot Interaction. *J. Hum.-Robot Interact.* 13, 2, Article 27 (June 2024), 29 pages. doi:10.1145/3659058
- [62] Azadeh Shariati, Mojtaba Shahab, Ali Meghdari, Ali Amoozandeh Nobaveh, Raman Rafatnejad, and Behrad Mozafari. 2018. Virtual Reality Social Robot Platform: A Case Study on Arash Social Robot. In *Social Robotics*, Shuzhi Sam Ge, John-John Cabibihan, Miguel A. Salichs, Elizabeth Broadbent, Hongsheng He, Alan R. Wagner, and Álvaro Castro-González (Eds.). Springer International Publishing, Cham, 551–560. doi:10.1007/978-3-030-05204-1_54
- [63] Sara Sheikholeslami, Ajung Moon, and Elizabeth A Croft. 2017. Cooperative gestures for industry: Exploring the efficacy of robot hand configurations in expression of instructional gestures for human-robot interaction. *The International Journal of Robotics Research* 36, 5-7 (June 2017), 699–720. doi:10.1177/0278364917709941
- [64] Masahiro Shiomi, Hidenobu Sumioka, and Hiroshi Ishiguro. 2020. Survey of social touch interaction between humans and robots. *Journal of Robotics and Mechatronics* 32, 1 (2020), 128–135. doi:10.20965/jrm.2020.p0128
- [65] Erica N. Shook, George Thomas Barlow, Daniela Garcia-Rosales, Connor J. Gibbons, and Tessa G. Montague. 2024. Dynamic skin behaviors in cephalopods. *Current Opinion in Neurobiology* 86 (June 2024), 102876. doi:10.1016/j.conb.2024.102876
- [66] Sichao Song and Seiji Yamada. 2018. Bioluminescence-Inspired Human-Robot Interaction: Designing Expressive Lights that Affect Human’s Willingness to Interact with a Robot. In *Proceedings of the 2018 ACM/IEEE International Conference on Human-Robot Interaction (Chicago, IL, USA) (HRI '18)*. Association for Computing Machinery, New York, NY, USA, 224–232. doi:10.1145/3171221.3171249
- [67] Maria Staudte and Matthew W. Crocker. 2009. Visual attention in spoken human-robot interaction (HRI '09). Association for Computing Machinery, New York, NY, USA, 77–84. doi:10.1145/1514095.1514111
- [68] Kyle Strabala, Min Kyung Lee, Anca Dragan, Jodi Forlizzi, and Siddhartha S. Srinivasa. 2012. Learning the communication of intent prior to physical collaboration. In *2012 IEEE RO-MAN: The 21st IEEE International Symposium on Robot and Human Interactive Communication*. 968–973. doi:10.1109/ROMAN.2012.6343875
- [69] Devi Stuart-Fox and Adnan Moussalli. 2008. Selection for Social Signalling Drives the Evolution of Chameleon Colour Change. *PLOS Biology* 6, 1 (Jan. 2008), e25. doi:10.1371/journal.pbio.0060025 Publisher: Public Library of Science.
- [70] Thomas Suddendorf, Gabrielle Simcock, and Mark Nielsen. 2007. Visual self-recognition in mirrors and live videos: Evidence for a developmental asynchrony. *Cognitive Development* 22, 2 (2007), 185–196. doi:10.1016/j.cogdev.2006.09.003
- [71] Daniel Szafir, Bilge Mutlu, and Terry Fong. 2015. Communicating Directionality in Flying Robots. In *Proceedings of the Tenth Annual ACM/IEEE International Conference on Human-Robot Interaction (Portland, Oregon, USA) (HRI '15)*. Association for Computing Machinery, New York, NY, USA, 19–26. doi:10.1145/2696454.2696475
- [72] Susumu Tachi, M Inami, and Y Uema. 2014. Augmented reality helps drivers see around blind spots. *IEEE Spectrum* 31 (2014).
- [73] Gilbert Tang, Phil Webb, and John Thrower. 2019. The development and evaluation of Robot Light Skin: A novel robot signalling system to improve communication in industrial human-robot collaboration. *Robotics and Computer-Integrated Manufacturing* 56 (2019), 85–94. doi:10.1016/j.rcim.2018.08.005
- [74] Georgios Tsamis, Georgios Chantziaras, Dimitrios Giakoumis, Ioannis Kostavelis, Andreas Kargakos, Athanasios Tsakiris, and Dimitrios Tzavaras. 2021. Intuitive and Safe Interaction in Multi-User Human Robot Collaboration Environments through Augmented Reality Displays. In *2021 30th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*. 520–526. doi:10.1109/RO-MAN50785.2021.9515474

- [75] Johan Wagemans, Joeri De Winter, Hans Op De Beeck, Annemie Ploeger, Tom Beckers, and Peter Vanroose. 2008. Identification of Everyday Objects on the Basis of Silhouette and Outline Versions. *Perception* 37, 2 (Feb. 2008), 207–244. doi:10.1068/p5825
- [76] Michael Walker, Hooman Hedayati, Jennifer Lee, and Daniel Szafrir. 2018. Communicating Robot Motion Intent with Augmented Reality. In *Proceedings of the 2018 ACM/IEEE International Conference on Human-Robot Interaction* (Chicago, IL, USA) (HRI '18). Association for Computing Machinery, New York, NY, USA, 316–324. doi:10.1145/3171221.3171253
- [77] Tim Wengefeld, Dominik Höchemer, Benjamin Lewandowski, Mona Köhler, Manuel Beer, and Horst-Michael Gross. 2020. A Laser Projection System for Robot Intention Communication and Human Robot Interaction. In *2020 29th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*. 259–265. doi:10.1109/RO-MAN47096.2020.9223517
- [78] Christian JAM Willemse, Alexander Toet, and Jan BF Van Erp. 2017. Affective and behavioral responses to robot-initiated social touch: toward understanding the opportunities and limitations of physical contact in human–robot interaction. *Frontiers in ICT* 4 (2017), 12. doi:10.3389/fict.2017.00012
- [79] Jacob O. Wobbrock, Leah Findlater, Darren Gergle, and James J. Higgins. 2011. The aligned rank transform for nonparametric factorial analyses using only anova procedures. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Vancouver, BC, Canada) (CHI '11). Association for Computing Machinery, New York, NY, USA, 143–146. doi:10.1145/1978942.1978963
- [80] J.D. Zamfirescu-Pereira, David Sirkin, David Goedicke, Ray LC, Natalie Friedman, Ilan Mandel, Nikolas Martelaro, and Wendy Ju. 2021. Fake It to Make It: Exploratory Prototyping in HRI. In *Companion of the 2021 ACM/IEEE International Conference on Human-Robot Interaction* (Boulder, CO, USA) (HRI '21 Companion). Association for Computing Machinery, New York, NY, USA, 19–28. doi:10.1145/3434074.3446909
- [81] Ran Zhao, Zhaopeng Zhu, Xiaotong He, Yue Jiang, Wenjie Wu, and Yun Wang. 2026. Ready for the Touch: Exploring Users' Perceived Transparency of Robot Pre-Touch Cues. In *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems* (CHI '26). Association for Computing Machinery, New York, NY, USA, Article 977, 16 pages. doi:10.1145/3772318.3791193
- [82] Caroline Yan Zheng, Yuting Chen, Adrian Latupeirissa, Georgios Andrikopoulos, Anna Ståhl, and Madeline Balaam. 2025. Towards Caring Touch From Technologies: Knowledge From Healthcare Practitioners. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems* (CHI '25). Association for Computing Machinery, New York, NY, USA, Article 505, 16 pages. doi:10.1145/3706598.3713736
- [83] Caroline Y. Zheng, Ker-Jiun Wang, Maitreyee Wairagkar, Mariana von Mohr, Erik Lintunen, and Aikaterini Fotopoulou. 2024. Simulating the psychological and neural effects of affective touch with soft robotics: an experimental study. *Frontiers in Robotics and AI* 11 (Nov. 2024). doi:10.3389/frobt.2024.1419262 Publisher: Frontiers.

A Performance Study

To give a better impression of the VR intent communication (cf., section 5), we provide example images of the robot conveying an invitation to be touched by the human, and its intention to touch the human, for BASELINE (gesture & gaze), PRE-MOTION MIRROR SKIN, and IN-MOTION MIRROR SKIN (Figure 6):

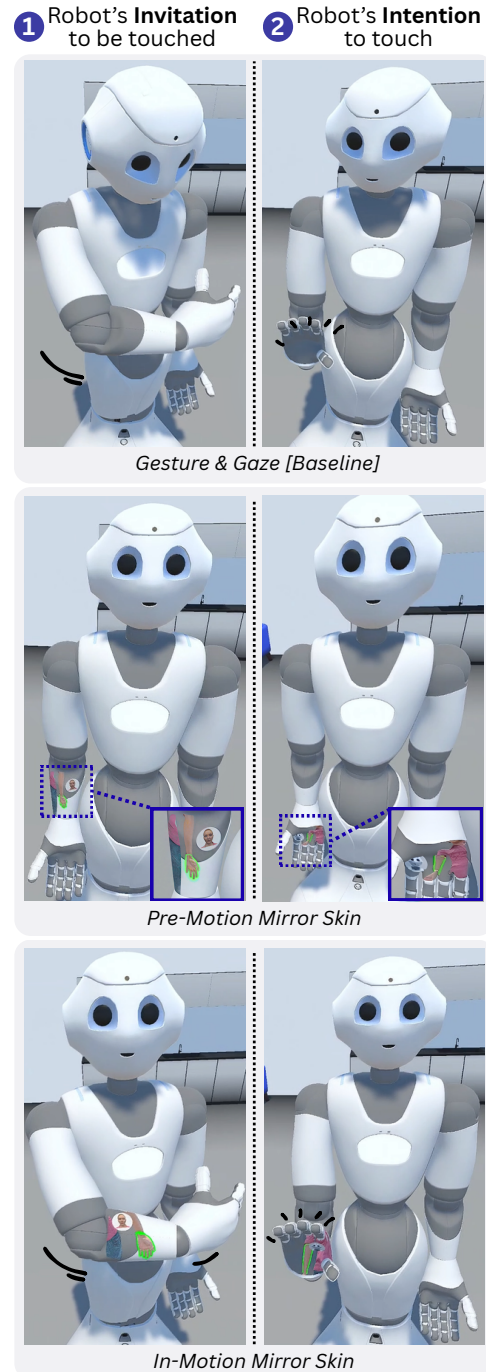


Figure 6: Examples of the robot communicating (1) an invitation to be touched, or (2) its intention to touch the human, for all three feedback techniques of study 2.