

Nonverbal Human Signals Can Help Autonomous Agents Infer Human Preferences for Their Behavior

Kate Candon
Jesse Chen
kate.candon@yale.edu
jesse.b.chen@yale.edu
Yale University
New Haven, CT, USA

Yoony Kim
Zoe Hsu
yoony.kim@yale.edu
zoe.hsu@yale.edu
Yale University
New Haven, CT, USA

Nathan Tsoi
Marynel Vázquez
nathan.tsoi@yale.edu
marynel.vazquez@yale.edu
Yale University
New Haven, CT, USA

ABSTRACT

An overarching goal of Artificial Intelligence (AI) is creating autonomous, social agents that help people. Two important challenges, though, are that different people prefer different assistance from agents and that preferences can change over time. Thus, helping behaviors should be tailored to how an individual feels during the interaction. We hypothesize that human nonverbal behavior can give clues about users' preferences for an agent's helping behaviors, augmenting an agent's ability to computationally predict such preferences with machine learning models. To investigate our hypothesis, we collected data from 194 participants via an online survey in which participants were recorded while playing a multiplayer game. We evaluated whether the inclusion of nonverbal human signals, as well as additional context (e.g., via game or personality information), led to improved prediction of user preferences between agent behaviors compared to explicitly provided survey responses. Our results suggest that nonverbal communication – a common type of human implicit feedback – can aid in understanding how people want computational agents to interact with them.

KEYWORDS

human-agent interaction; nonverbal behavior; preference learning

ACM Reference Format:

Kate Candon, Jesse Chen, Yoony Kim, Zoe Hsu, Nathan Tsoi, and Marynel Vázquez. 2023. Nonverbal Human Signals Can Help Autonomous Agents Infer Human Preferences for Their Behavior. In *Proc. of the 22nd International Conference on Autonomous Agents and Multiagent Systems (AAMAS 2023)*, London, United Kingdom, May 29 – June 2, 2023, IFAAMAS, 10 pages.

1 INTRODUCTION

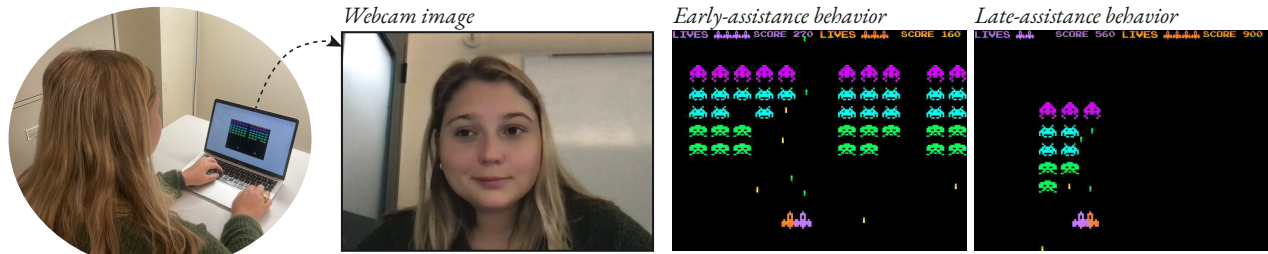
An overarching goal of Artificial Intelligence (AI) is creating autonomous, social agents that help people [10, 14, 25, 43, 68]; yet, human-agent interactions are often scripted and rigid, lacking the adaptability typical of human social encounters. One common approach to address this limitation is through eliciting explicit preferences (e.g., [23, 63]) because different people prefer different assistance from autonomous agents. In the future, it would be advantageous if autonomous agents understood some of the factors influencing user preferences such that they could not only adapt in a post-hoc manner but more proactively change their behavior to fit users' desires. Even more, it would be useful if autonomous

agents could understand implicit social cues provided by people during interactions. Information inferred from these cues could augment explicitly provided feedback, improving an agent's ability to model human preferences and adapt their behavior accordingly.

In contrast to existing autonomous agents, people are often able to adapt their behavior to other people by interpreting social cues provided in human-human interactions. For example, expert teachers are able to recognize students' affective states so they can then adjust the pace and content of the learning material [49]. Elder adults infer which caregivers are available to provide assistance based on body orientation, head position, and gaze [74]. Individuals select acceptable partners for different social goals based on observable nonverbal behavior [66]. But in human-agent interaction scenarios, it is unclear whether humans would provide useful nonverbal cues in response to an agent's actions. People may behave differently when interacting with an autonomous agent than when interacting with another person [2, 18, 20, 21, 28, 53, 58]. For instance, people have been found to provide fewer nonverbal signals when interacting with a robot if there are no other humans present [29] and show fewer physiological signs of emotion in response to agents compared to humans [62, 70].

Another challenge in creating autonomous agents that reason about implicit social cues from people is that these cues can be hard to understand at times. For example, when a person is engaged in solving a task with an agent, are human social cues due to the agent's behavior or due to the person's own actions in the task? Also, social cues may have different meanings in different situations. Smiles, for instance, are often considered an expression of enjoyment [27]. Nonetheless, people have been shown to smile in response to robot mistakes [29], for which we may naturally expect expressions of surprise, disappointment or disapproval instead. Even if people do emit nonverbal signals when collaborating with an autonomous agent, will the signals be useful to the agent?

Despite the above challenges, we hypothesize that nonverbal human signals can provide clues about how people want computational agents to interact with them. To investigate this idea, we conducted an online study in which participants interacted with an autonomous agent in a fast-paced collaborative task. As shown in Figure 1, the participants played two games of a multi-player video game with different agent behaviors. We did not prime participants for cooperation (nor competition) in the game because we wanted to see how they would naturally react to an agent that tried to help them in the game when this help was not necessarily expected a priori. Also, we recorded participants via their webcams during interactions to later analyze their nonverbal signals.



(a) Experimental setup: Participants played an online, multi-player version of Space Invaders. Their faces were recorded while playing the game.

(b) Participants (purple spaceship) experienced two types of helping behaviors by the co-player (orange spaceship) in the study.

Figure 1: We collected data from an online human-agent interaction to investigate the usefulness of including human nonverbal signals into models for predicting which agent behavior participants preferred. In our study, the participants were recorded while playing the game with an autonomous agent via their webcams.

We present results from multiple analyses to understand how different sources of information impact the prediction of user preferences for agent behaviors in the multi-player game scenario. First, we study the possibility of predicting preferences based on impressions of the interaction reported via surveys. These types of subjective impressions are commonly gathered in user studies [5, 47, 50] and survey data can help understand key factors that influence preferences. Then, we investigate whether the inclusion of nonverbal human signals improves preference predictions, per our hypothesis. Finally, we investigate whether additional context information, like information about the state of the game, further helps infer preferences over agent behaviors. To the best of our knowledge, this is the first study to investigate the usefulness of nonverbal human signals in predicting user preferences in a fast-paced and collaborative human-agent interaction scenario. The anonymized data that we used for our analyses can be found in https://github.com/yale-img/collabHAI_pref to facilitate future replication efforts and more complex preference modeling.

Overall, our findings provide insights for creating real-time adaptable autonomous agent behavior that leverages spontaneous nonverbal human reactions in the future. This could potentially help reduce the need to query users often in order to understand which agent behaviors work well in comparison to others.

2 RELATED WORK

Implicit human feedback encompasses a wide range of spontaneous behavior and reactions humans provide during interactions [65]. Understanding implicit human cues has been a longstanding goal in affective computing [26, 57]. Researchers have leveraged nonverbal human signals to adapt the behavior of in-home devices [15, 75], improve generative deep learning models [40], and better support virtual negotiations [71]. Gaze has also been studied as a replacement for a “wake word” for a smart speaker [54] or to provide clues about a person’s intent [1], and facial and bodily expressions have been used to measure engagement in game tasks [31].

Work in affective computing has used nonverbal behavior to adapt agent behavior, but is typically focused on specific affective states or event detection. For example, Guerdan et al. [35] use human responses to explanations provided by an agent, focusing on specific affective states such as perceived challenge or competence.

Leite et al. [48] proposed an empathetic social robot that behaved differently based on perceived user boredom, interest, or frustration. Also, nonverbal human signals have been used to detect agent errors [44, 67, 72]. Unlike the previous corpus of work, we are not focused on specific emotion or event recognition. Rather, our work investigates if nonverbal reactions provide insight into a user’s preferences for helping behaviors by an autonomous agent.

We acknowledge that there are limitations to interpreting facial expressions, such as different representations of emotions across different cultures [4, 39]. However, we posit there is useful information about participant preferences to be gleaned from analysis of nonverbal behavior. Nonverbal signals could provide feedback to an agent without burdening participants for explicit feedback [52], e.g., via a think-aloud protocol.

Our motivation for studying nonverbal human reactions during human-agent interactions stems from recent work on adapting agent behavior based on implicit feedback with reinforcement learning [17, 51]. Different to these prior efforts, though, we do not bias humans’ internal goals or rewards in our study. Nor do we ask our participants to be expressive in a particular manner during human-agent interactions [51, 75] because this can result in fatigue. Our goal is to instead understand if their natural reactions while playing a collaborative game with an agent are useful for predicting user preferences for agent behaviors.

Research has shown that agents can learn to adjust their behavior based on explicit user preferences [7, 22, 32, 42, 61]. Querying users for preferences involves interaction costs, so previous work has investigated how to ask for preferences in a way that minimizes annoyance [34] and reduces the number of queries needed for adaptation [9, 73]. To complement this line of work, we investigate:

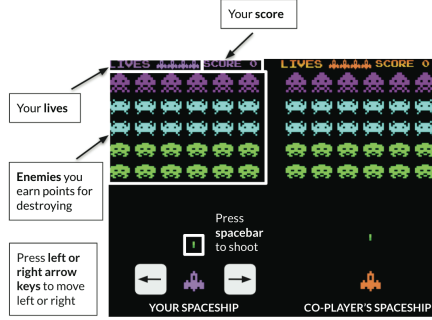
RQ1: *Can natural, nonverbal human reactions be leveraged to better predict user preferences for an agent’s helping behaviors compared to only using explicit survey responses about game and agent attributes?*

Because recent work in psychology suggests that context is key when interpreting nonverbal human behavior [4, 39], we also ask:

RQ2: *Does the inclusion of additional context, beyond human nonverbal reactions, help predict human preferences over agent behaviors?*

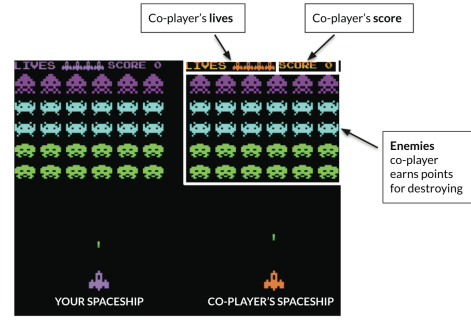
Following Schilit and Theimer [64], we broadly consider context to be any element providing information about the environment

You are the **purple ship**, and you start with 4 lives. Use the **left and right arrow keys** to move, and press the **spacebar** to shoot. Your shooting speed is limited, so pressing the spacebar may not always result in a bullet being shot—that's part of the game. **You get points for enemies destroyed on the left half of the screen.**



(a) Introduction to the participant player

There will be an **orange co-player** in the game, automatically controlled by **artificial intelligence**. The co-player gets points for enemies destroyed on the right half of the screen.



(b) Introduction to the co-player (with AI identity)

Figure 2: Participants were introduced to the game as shown above. The explanation included explicit pointers to the scores and number of remaining lives for each player, where the enemies would appear, and how to move the spaceship and shoot bullets.

in which an agent is situated. Nonverbal human reactions can be considered as context for an agent in our multi-player game. Furthermore, the background and personality traits of users as well as task (or activity) statistics are additional contextual factors relevant to the human-agent interactions studied in this work.

3 METHOD

In order to investigate the research questions outlined in the prior Section, we collected data through an exploratory online interaction. For the interaction, the participants completed a web survey through which they interacted with an agent in a two-player version of the Space Invaders game (Figure 1). The participant controlled one spaceship and the other spaceship, referred to as the *co-player*, was algorithmically controlled using simple heuristic behaviors. With their consent, we recorded participants during the games of Space Invaders via their webcams. The online study that we conducted for this data collection was approved by our local Institutional Review Board. We used the collected data to build models that predicted which co-player behavior was preferred by the participants, which is the focus of this paper. For additional motivations, details, and results of the user study, please refer to the work by Candon et al. [12].

3.1 Data Collection

3.1.1 Participants. We recruited 360 participants for the study through Prolific [55]. The recruitment criteria required participants to be 18 years of age or older, be fluent in English, reside in the United States, and have normal or corrected-to-normal vision.

Out of the 360 participants, 194 participants were included in the final dataset, and 166 participants were excluded. Participants were excluded if they asked to withdraw, if their game logs indicated poor Internet connection, or if there was insufficient webcam data captured during the game. More details about our participant exclusion criteria are in Section A of the Supplementary Material.

Of the 194 final participants, 78% identified as female, 20% identified as male, 1% identified as nonbinary, and 1% preferred not to

say. The participants' ages ranged from 18 to 66 years old, with an average age of 26.87 years ($SD = 9.30$). Participants indicated using computers often: 93% of participants used a computer daily, 6% used a computer 4-6 times a week, and 1% used a computer 2-3 times a week. A little more than half of the participants (53%) played video games once a week. When asked specifically about Space Invaders, 45% of the participants reported that they had played the game before, 46% reported they had not played it, and 8% were unsure.

3.1.2 Space Invaders Game. The Space Invaders game was introduced to participants as shown in Figure 2. The participant controlled a purple spaceship and the participant's co-player was an orange spaceship. The participant's spaceship started on the left side of the screen, and the co-player's spaceship started on the right side of the screen, but both could move left and right within the full bounds of the game screen. Both spaceships could shoot upwards to destroy enemies, and each was assigned points individually for enemies destroyed on the side of the screen on which they originally started, regardless of whose bullet destroyed the enemy. Both the participant and co-player started the game with four lives. A player lost a life when hit by an enemy or a bullet from the enemies. Enemies moved left and right across the screen and slowly downwards, closer to the spaceships, until they were hit by a bullet or reached the bottom of the game screen. A game over screen appeared showing the final scores once all enemies were destroyed, both players lost all their lives, or an enemy reached the bottom of the game screen.

3.1.3 Co-player Behaviors. Because we were interested in predicting preferences between agent behaviors, we created two different co-player behaviors: early-assistance and late-assistance. In our game, the co-player could provide assistance to the participant by travelling to the participant's side of the game screen. Once on the participant's side, the co-player could help destroy enemies for which the co-player received points. We designed the two behaviors so they differed by the timing of when the co-player travelled to the participant's side of the screen to provide assistance. In the *early-assistance* behavior, the co-player went over to the participant's

side of the screen on two occasions during the game while there were still enemies on the co-player’s side of the screen (Figure 1b, left). In the *late-assistance* behavior, the co-player only when to the participant’s side of the screen after all of the enemies on the co-player’s side were already destroyed (Figure 1b, right).

3.1.4 Procedure. For the interaction, participants completed an online Qualtrics survey, which included two games of Space Invaders. The participant first consented to participate in the study, to have their video recorded, and to have their images and video clips shared publicly. The survey then included a video check to ensure that their webcam was working and that their face could be detected in the webcam images. The participants then completed a demographic section of the survey, as discussed in Section 3.1.1. This section also included personality data via the Revised Competitiveness Index [38] and the Ten Item Personality Measure (TIPI) [30]. The survey then introduced the Space Invaders game with a combination of text explanations and visual instructions, as illustrated in Figure 2. We purposefully did not prime participants for cooperation nor competition with the co-player. The participants experienced two games of Space Invaders, each followed by a post-game survey about their perceptions of the game and of the co-player. Each game involved a different co-player behavior, as described in Section 3.1.3. Participants were not informed of the order in which they experienced the two behaviors, and the order of behaviors was counterbalanced between participants. After playing two games of Space Invaders and answering both sets of post-game questions, participants were asked a final set of questions about the differences between the games and their preferences. Participants were presented with an optional Berkeley Expressivity Questionnaire (BEQ) [33]. The study took about 18 minutes to complete. Individuals were paid \$3.60 for participating.

3.2 Preference Prediction Task

In this work, we study the problem of predicting which co-player behavior a participant reported that they preferred at the end of their session. We consider this problem a multi-class classification problem because the final section of our survey asked the participants to state whether they preferred the first co-player behavior, the second one, or did not prefer one over the other. Based on which order the participant experienced the two behaviors, we encoded the targets as either *Early*, *Late*, or *No Preference*.

3.2.1 Inputs to Preference Classifiers. We considered different combinations of four types of input features for preference classifiers. These feature types corresponded to 1) features derived from post-game survey responses, 2) nonverbal reaction data, 3) participant demographic data, and 4) game context. These four types of data were selected for different reasons. First, survey data is commonly used for understanding human preferences over agent behaviors (e.g., [16, 24, 69]). Second, nonverbal reaction data corresponded to implicit feedback that we hoped could be leveraged to better understand human preferences (per *RQ1*). Third, we suspected individual human factors would affect preferences, so we included demographic data (per *RQ2*). Finally, because our interaction domain was dynamic, it was important to consider game data to understand what the user and agent were doing and the state of

the environment (per *RQ2*). While it is rare to consider all four data types in conjunction, we believe that there is value in systematically studying how they can all be used to predict user preferences.

For our analysis, we considered a variety of preference classifier models that differed in terms of the features they received as input. First, we considered models that took as input survey features only (Survey). These models were regarded as baselines because survey responses are commonly used to understand human perceptions of an interaction (e.g., [5, 47, 50]). Second, we added nonverbal reaction data (Survey+Nonverbal). Comparing Survey and Survey+Nonverbal allowed us to experimentally investigate whether nonverbal data helped predict user preferences (*RQ1*). Third, we considered models that also took as input demographics and game data, to incorporate additional context. This led to three more sets of inputs: Survey+Nonverbal+Demo, Survey+Nonverbal+Game, and Survey+Nonverbal+Demo/Game. These feature combinations allowed us to experimentally investigate whether additional context further helped predict user preferences (*RQ2*). The next sections provide more detail about each specific feature type.

Post-Game Survey Responses: Our first set of input data was encoded from survey questions participants were asked after playing each game of Space Invaders. The questions were about the game experience, the perception of the co-player, and whether or not the participant thought they had helped the co-player. Game experience questions included agreement with statements from Large et al. [46] about whether the game was enjoyable, difficult, boring, and fun. Co-player perception questions included agreement with statements from Large et al. [46] about whether the co-player was helpful, proficient, intelligent, and annoying. Co-player perception questions also included if the participant liked the behavior of the co-player or if they thought anything about the behavior of the co-player seemed unusual. In addition, the participants evaluated the level of competence, warmth, and discomfort of the co-player using the 18 attributes from the Robotic Social Attributes Scale (RoSAS) [13]. More details about survey questions are included in Section B.1 of the Supplementary Material.

We considered both a Full set of survey features and a Selected set of survey features, which we experimentally found to be most important for preference prediction. For the Full set of survey features, we processed the raw data (e.g., via scaling and one-hot encoding) for 14 survey questions to arrive at 18 features encoding survey information for each game (as detailed in Section B.1 of the Supplementary Material). We additionally explored reducing the number of included survey features using the notion of Gini importance [11] from a trained Random Forest model, as implemented by the scikit-learn Python Library [56]. In particular, we selected the six post-game survey features with importance greater than 0.05: liked behavior, competence, annoyingness, helpfulness, warmth, and discomfort. Additionally, both sets of survey features included a feature encoding the order in which the participant experienced the two games, as well as a one-hot encoding for which of three co-player identities the participant experienced. We provided the survey data to a classifier in two ways: passing features for both games by concatenating them (Concat), and passing the difference between the value from the game with the early-assistance behavior and the game with the late-assistance behavior (Diff).

Nonverbal Reactions: We analyzed human facial and body reactions captured while the participants played Space Invaders. Our version of the Space Invaders game captured images of a participant via the participant’s own webcam at a framerate of 15 frames per second. We analyzed the images automatically using OpenFace 2.0 [3], a open-source toolkit for automatic behavior analysis. For each image, OpenFace 2.0 [3] extracted information about head pose, eye gaze, facial landmarks, and facial action units.

We explored different ways to incorporate nonverbal reaction data. Table 1 describes the features and the summary statistics that we used to arrive at 59 nonverbal reaction features for each game. For this set of features, we considered four approaches to provide them to a classifier: Full, Visit-Split, Visit-Diff, and Visit-Hadamard. In the Full case, we concatenated the summary features for both games and input them into the model. In the other cases, we computed summary features for when the co-player visited the participant’s side of the screen (“co-player visit”), and when it was on its side (“no co-player visit”).¹ In the Visit-Split case, both subsets of summary statistics for both games were concatenated and input to the model. In the Visit-Diff case, we computed the difference between the “co-player visit” and “no co-player visit” statistics within each game, concatenated the results, and then input them into to the model. Lastly, for Visit-Hadamard, we computed the element-wise product between the “co-player visit” and “no co-player visit” statistics within each game, concatenated the results, and passed the output to the classifier. Because the Nonverbal feature sets were much larger than the other feature sets, we considered a reduced set of nonverbal features derived by applying Principal Component Analysis (PCA) with 5, 10, 20, and 40 principal components in addition to the full set of original features.

Demographics: We preprocessed raw self-reported demographics information from the survey to arrive at 20 features describing the personality, age and gender of each participant as well as how often this person played video games, used a computer, and whether they had played Space Invaders before (as detailed in Section B.2 of the Supplementary Material). Similar to the survey data, we considered a Selected set of five demographic features, again with Gini importance from a trained RF model. The Selected set of demographic features included: how often the participant played video games, competitiveness index [38], negative expressivity from BEQ [33], and the extroversion and emotional stability dimensions of TIPI [30].

¹We grouped data based on the location of the co-player because we suspected that participants would react differently when it was on their side of the game screen. Also, in order to capture facial reactions as the co-player was traveling to or back from the participant’s side, we included 100 frames before and after a co-player visit when splitting the facial features into the “co-player visit” and “no co-player visit” sets.

Game Context: To enable classifier models to reason about contextual factors related to what occurred during Space Invaders, we included information from game logs. These game logs contained information about the state of the game, participant actions, and co-player actions for each rendered frame of Space Invaders. In particular, we extracted and analyzed: number of times the co-player and the participant moved to opposite sides of the game screen, the total number of frames in which the co-player and the participant stayed on opposite sides of the game screen, the number of participant enemies destroyed by the co-player, the number of co-player enemies destroyed by the participant, and the participant and co-player’s final scores and number of lives remaining. We provided the game data to a classifier in two ways: passing features for both games by concatenating them (Concat), and passing the difference between features for the two games (Diff).

3.2.2 Models for Predicting Participant Preferences. We considered various popular Machine Learning (ML) algorithms from the scikit-learn Python library [56]: Support Vector Machines (SVM), Random Forests (RF), K-Nearest Neighbors (KNN), and Multi-Layer Perceptron (MLP). We selected these simple, well-established classifiers because more complex models are likely to overfit on our small dataset and also tend to be more opaque. For each classifier, we performed a gridsearch over a range of suitable hyperparameters with survey data as inputs. We then used those hyperparameters for all models trained on different permutations of input feature sets. Details of the gridsearch procedure are presented in Section C of the Supplementary Material.

3.2.3 Evaluation of Preference Classifier Performance. We evaluated the classifiers using $F_1\text{-Score} = 2(\text{precision}^{-1} + \text{recall}^{-1})^{-1}$ because it balances different kinds of prediction errors and is less biased by class imbalance [36]. $\text{Precision} = |TP|/(|TP| + |FP|)$ is the proportion of positive predictions that are true positive targets, where $|TP|$ is the number of true positives and $|FP|$ is the number of false positives. $\text{Recall} = |TP|/(|TP| + |FN|)$ is the proportion of true positive targets that are correctly identified, where $|TP|$ is the number of true positives and $|FN|$ is the number of false negatives. We use the macro notion of $F_1\text{-Score}$. That is, we calculate the $F_1\text{-Score}$ for each of the three target classes and take the average.

In order to make the most use of our limited number of samples, we evaluated models using leave-one-out cross validation (LOOCV), rather than dividing the dataset into static training and testing sets. For each participant, we trained a model on the data from the other 193 participants and made a single test prediction. We then calculated $F_1\text{-Score}$ from the confusion matrix created from the 194 individual predictions (i.e., from the 194 folds of LOOCV).

Table 1: Description of Open Face 2.0 [3] attributes used and which summary statistics (standard deviation (stdev), mean, and/or maximum value) from the capture frames were included.

Attribute	Description	Summary Statistics
Gaze	Averaged gaze direction for both eyes in left-right and up-down directions	stdev
Translational Pose	Location of head with respect to camera	stdev
Rotational Pose	Pitch, yaw, roll for head	stdev
Facial Action Unit Intensity	Intensities (from 0 to 5) of 17 action units	stdev, mean, maximum

4 RESULTS

This section presents our evaluation of ML algorithms to predict which co-player behavior the participants preferred. We first discuss three naive baselines for predicting preferences. Then, we describe our analysis on whether or not nonverbal signals aid in modeling co-player behavior preferences (RQ1). Finally, we discuss results about incorporating additional context into our models (RQ2) and investigate feature importance values to gain insights into what kind of information is being leveraged by our preference classifiers.

4.1 Naive Baselines for Predicting Preferences

Of our 194 participants, 116 (60%) preferred the late-assistance co-player behavior, 50 (26%) preferred the early-assistance co-player behavior, and 28 (14%) did not prefer one over the other. With this distribution, always selecting the most dominant class results in an F_1 -Score of 0.25. Randomly sampling from the three classes and running the sampling $N = 10$ times results in an F_1 -score of $M = 0.29$ ($SD = 0.02$). Weighted sampling from the distribution of the true labels results in an F_1 -Score of $M = 0.33$ ($STD = 0.03$, $N = 10$).

4.2 Nonverbal Reactions Can Help Model Preferences

We first investigated whether nonverbal human reactions can help predict user preferences for an agent’s helping behaviors compared to only using explicit survey responses (RQ1). Table 2 shows the highest F_1 -Scores for the ML algorithms that were trained on both the Survey and Survey+Nonverbal input permutations described in Section 3.2.1. All F_1 -Scores are notably higher than the F_1 -Scores from naive baselines in Section 4.1. For each of the SVM, RF, and MLP algorithms, the highest F_1 -Score was from the set of models that incorporated nonverbal information. For KNN, the Survey only inputs resulted in a higher F_1 -Score than Survey+Nonverbal inputs. However, KNN was the lowest performing classifier among the ML algorithms considered in this work. We suspect that the low performance of KNN is related to the algorithm’s known trouble with the presence of outliers in training data [8] since nonverbal reactions can vary greatly. KNN’s low performance could also be due to the increase in dimensionality of input features with the addition of nonverbal data given our limited-sized dataset.

4.3 Incorporating Additional Context

Driven by the question of whether additional context information can help predict human preferences (RQ2), we analyzed the performance of preference classifiers with additional permutations of input data (Section 3.2.1). In particular, we considered demographic information, game context, and the combination of the two as additional context that could help reason about human perceptions of the agent. Figure 3 presents the highest F_1 -Score within each of the five permutations of input data for each ML algorithm considered in this analysis. Table 1 of the Supplementary Material includes information about the input features that resulted in the highest F_1 -Score for each type of classifier.

For each type of classifier, the highest F_1 -Score was from a set of models that included nonverbal reaction data and at least one other kind of additional context. Across all classifiers and all input

Table 2: Best F_1 -Scores for each classifier for Survey and Survey+Nonverbal input combinations. F_1 -Score was calculated over a confusion matrix derived from individual predictions of 194 folds of LOOCV.

Classifier	Survey	Survey+Nonverbal
Support Vector Machine	0.50	0.56
Random Forest	0.52	0.55
Multi-Layer Perceptron	0.51	0.54
K-Nearest Neighbors	0.50	0.47

permutations, SVM with Survey+Nonverbal+Game had the highest F_1 -Score (0.60). For both RF and MLP, the set of models with Survey+Nonverbal+Demographics/Game inputs had the highest F_1 -Scores of 0.59 and 0.57, respectively. The highest F_1 -Score for KNN (0.54) was obtained with the Survey+Nonverbal+Demographics data. Once again, KNN underperformed the other classifiers. Again, all F_1 -Scores are notably higher than the F_1 -Scores from naive baselines in Section 4.1.

Notably, including both kinds of additional contextual information did not always result in the highest F_1 -Score (see results for SVM and KNN in Figure 3). This means that one cannot take for granted that including more features will result in higher performance. Rather, it is important to further explore how to incorporate additional context when reasoning about internal human states.

In addition to classifier performance, we were interested in understanding which features played a key role in predicting participant preferences for agent helping behaviors. First, we considered doing this analysis with the SVM model because it had the top performance; but it employed a radial basis function kernel, making it hard to disentangle feature importances. Thus, we instead analyzed feature importance values with the RF algorithm. The RF models had good performance. Also, the high-level of interpretability of the underlying decision trees used for the preference predictions made it easy to understand what features mattered.

We compared the importance of features for a model considering Survey+Nonverbal inputs and Survey+Nonverbal+Demo/Game inputs. Instead of using LOOCV as for the prior results, we trained one RF model for each input combination with the data from all 194 participants for this analysis. This resulted in two models for which we then calculated Gini feature importance values [11]. Gini importance indicates the total decrease in node impurity, weighted by the probability of reaching that node in the tree, averaged over all trees of the forest. Intuitively, this importance value can be thought of as a measure of a features’ contribution to the homogeneity of the nodes and leaves in the forest.

Figure 4 illustrates the feature importance across the four types of input features considered by the RF with Survey+Nonverbal and Survey+Nonverbal+Demo/Game input combinations. In both cases, survey features were the most important to classifiers. Additionally, the results confirm our prior findings suggesting that the inclusion of additional context, via demographic and game information, contributes to improving classifier performance.

To investigate individual features, we examined feature importance with the Survey+Nonverbal+Demo/Game model. The three most important features were from survey responses: how much the

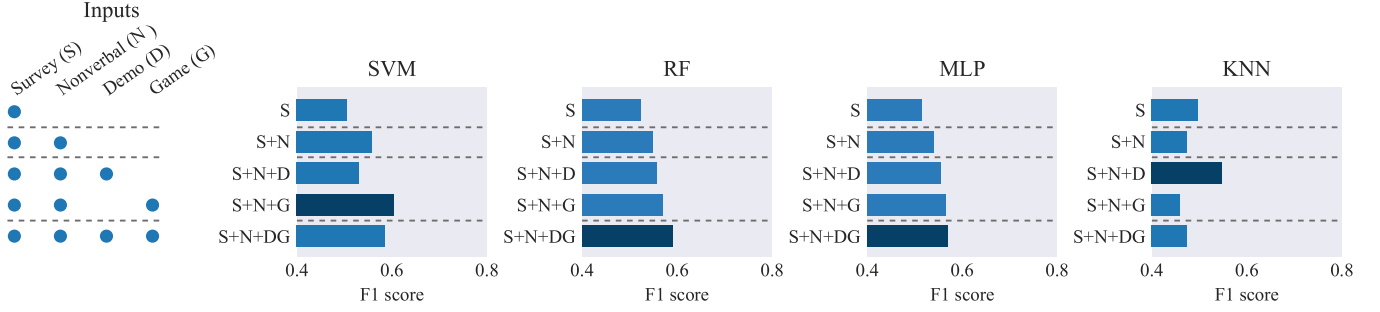


Figure 3: Best F_1 -Scores for each input combination. F_1 -Score was calculated over a confusion matrix derived from individual predictions of 194 folds of LOOCV. Dots on left indicate which information was considered in the model. Machine learning algorithms (SVM, RF, MLP, KNN) are ordered left-to-right in decreasing order of highest F_1 -Score. The darkest bar highlights the highest F_1 -Score for each algorithm. Dotted lines separate input combinations into: Survey only (original baseline), Survey + Nonverbal (new baseline), Survey + Nonverbal + one type of additional interaction context, and Survey + Nonverbal + both types of additional interaction context.

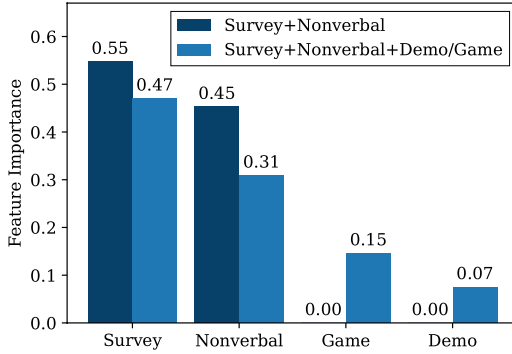


Figure 4: Feature importance by type of input feature.

participant liked the behavior of the agent (1st: 0.16), how annoying the participant found the agent (2nd: 0.15), and the competence of the agent (3rd: 0.08). The fourth (0.08) and fifth (0.05) most important features were principal components from the PCA analysis of nonverbal features. These principal components were driven by features describing AU45 (the “blink” action unit) for the game with the late-assistance co-player behavior as well as features describing AU2 (“outer brow raiser”), AU15 (“lip corner depressor”), and AU20 (“lip stretcher”) in the game with the early-assistance co-player behavior. The most important game feature was the number of frames the co-player was on the left side of the game screen (7th: 0.04). The most important demographic feature was how often the participant reported that they played video games (9th: 0.03).

5 DISCUSSION

Our results support the idea that implicit feedback in the form of nonverbal behavior can help model user preferences. We trained a variety of ML classifiers to predict co-player behavior preferences and evaluated them based on F_1 -Score. Across four different types of ML algorithms, the three with the highest F_1 -Score all had better performance when implicit nonverbal clues were included as inputs in addition to explicitly provided survey responses. All four ML

algorithms outperformed naive baselines based on the distribution of true labels. Repeatedly asking participants for their preferences can be annoying [34] or cause participants to lose interest in the interaction [59], so it is encouraging that analyzing nonverbal human signals can provide information about their preferences.

Figure 5 provides illustrative examples to highlight the value of implicit feedback, and why including it might help us better understand participant preferences. Participant P0354Z preferred the late-assistive behavior, which they experienced after the early-assistive behavior. This participant was very expressive as can be seen in Figure 5(a) and 5(b). They moved their eyebrows, changed the expression of their mouth, and opened their eyes wide. A similar image highlighting implicit feedback provided by another, less expressive participant is included in Section F of the Supplementary Material. Additional examples are included in our supplementary video. We are excited about the potential for computational agents to leverage this kind of information to better render prosocial behavior and proactively cooperate with users in the future.

Additionally, our results suggest that considering additional context (participant and game information in our case) is important to interpret nonverbal behavior. For all four ML algorithms, the set of models with the highest F_1 -Score not only included nonverbal data, but also included additional context via demographic data, game data, or both. This aligns with research in psychology that contests the assumption that emotions are recognized and communicated universally with particular facial expressions and argues for the importance of considering the context of interactions [4, 39].

At first it may appear obvious that more features would lead to better predictions, but that is not necessarily the case. Jensen and Shen [41] argue against the notion that more features in datasets translates to better performance due to feature redundancy and relevancy [19, 45] as well as the curse of dimensionality [6]. Thus, it is important to carefully study what information predictive models should include when an autonomous agent is reasoning about how favorably a user is viewing their behavior.

Being able to predict user preferences over agent behaviors would open up doors for agents to reason about which behaviors are “better” for a given user and, thus, incorporate more of those

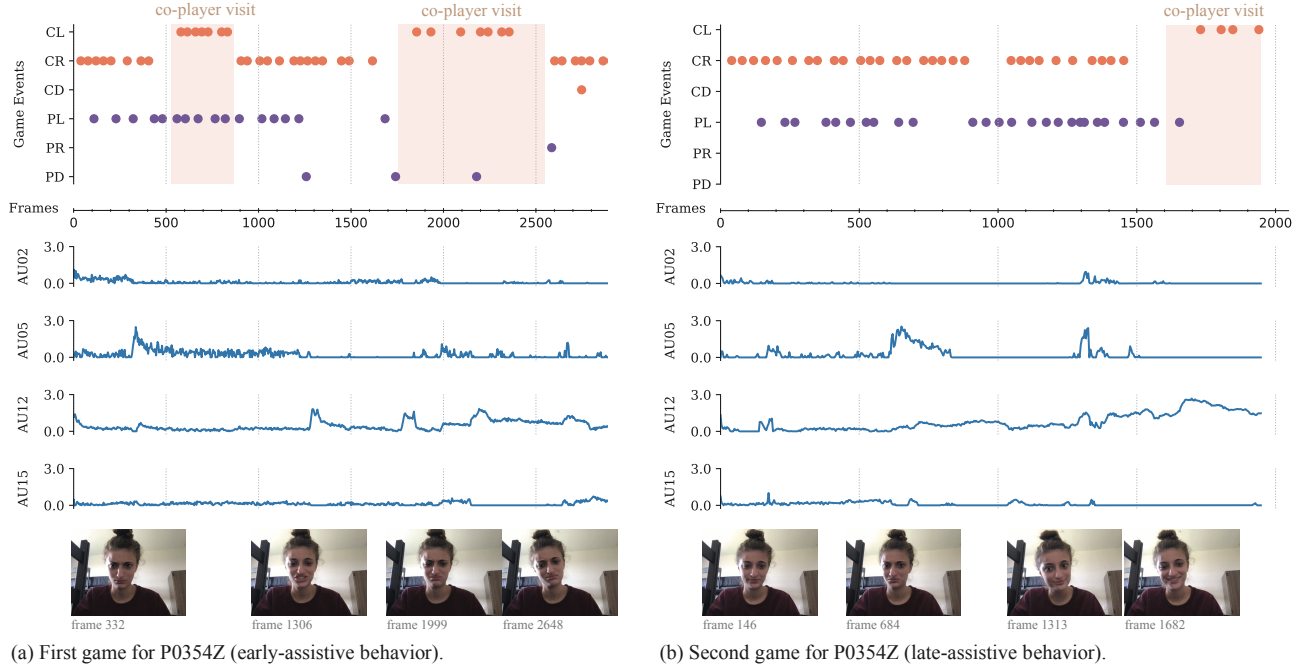


Figure 5: Visualization of data for participant P0354Z. Each plot shows game data (top) and webcam data (bottom). Game data includes six game events: co-player destroys enemy on the left side of the game screen (CL) or right side (CR), co-player dies (CD), participant destroys enemy on the left side (PL) or right side (PR), and participant dies (PD). The orange highlight identifies when the co-player is on the participants’ side (left side). Webcam data includes predicted action units (AU) by OpenFace: AU02 (outer brow raiser), AU05 (upper eyelid raiser), AU12 (lip corner puller), and AU15 (lip corner depressor). The x-axis corresponds to frame number as the game progresses from start to finish. All participants consented to having their images shared publicly.

behaviors into their interactions. Going forward, a better understanding of implicit human signals could also help human-agent interactions in other ways, e.g., enabling agents to learn how and when to ask a human collaborator for help [60].

6 LIMITATIONS AND FUTURE WORK

Our work has limitations, which motivate interesting future research directions. First, our analysis is bound to the domain of Space Invaders. It would be interesting to investigate predicting preferences using human nonverbal reactions in other interactive scenarios, including human-agent interactions that occur in person (e.g., with embodied virtual agents or robots).

Second, future work could explore a richer utilization of implicit feedback. We found encouraging results from simple ML models using summary statistics of nonverbal reaction features over frames of a game, but stronger results may be discovered if we consider the temporal nature of implicit feedback with more powerful ML algorithms (such as recurrent neural network models [37]).

Lastly, while this line of work is exciting, we must be cognisant of the ethical implications of designing autonomous agents able to analyze our nonverbal behavior and make inferences about our preferences. Going forward, it will be important to respect individual privacy and ensure individuals interacting with such autonomous agents are aware of the capabilities of the agents. Limitations on

the social manipulation of autonomous agents will also be critical as agents become better able to understand humans in interactions.

7 CONCLUSION

This work investigated the usefulness of natural nonverbal human signals to predict preferences between two agent behaviors. We collected data via an online interaction in which participants played two games of Space Invaders with an autonomous co-player exhibiting different behaviors. We built models to predict which of the two games was preferred by the participant and analyzed results from different combinations of input data for the models. Without biasing humans to be expressive, we found that we could leverage “free” information that they provided via nonverbal reactions to improve our ability to predict their preferences for agent behaviors.

Based on our findings, we propose two key recommendations for future interactive agents. First, we recommend designing agents with the capability to reason about nonverbal human reactions. This capability can improve the speed at which interactive agents adapt to personal preferences because nonverbal signals are readily available during interactions. Second, it is important to incorporate additional types of context, such as user personality or task statistics, into models that interpret nonverbal human signals. In our future work, we plan to take advantage of these insights to design better cooperative agents that reason about and adapt to the preferences of users with whom they interact.

ACKNOWLEDGMENTS

Thank you to the National Science Foundation (IIS-2106690) for partially supporting this work. Any opinions, findings, and conclusions or recommendations expressed in this paper are those of the authors and do not necessarily reflect the views of the NSF. Z. Hsu, Y. Kim, and J. Chen were partially supported by the Yale STARS program, the Yale College Dean's Research Fellowship, and the Yale Summer Experience Award.

REFERENCES

- [1] Henny Admoni and Siddhartha Srinivasa. 2016. Predicting user intent through eye gaze for shared autonomy. In *2016 AAAI Fall Symposium Series*.
- [2] Elisabeth André and Catherine Pelachaud. 2010. Interacting with embodied conversational agents. In *Speech technology*. Springer, 123–149.
- [3] Tadas Baltrušaitis, Amir Zadeh, Yao Chong Lim, and Louis-Philippe Morency. 2018. Openface 2.0: Facial behavior analysis toolkit. In *2018 13th IEEE international conference on automatic face & gesture recognition (FG 2018)*. IEEE, 59–66.
- [4] Lisa Feldman Barrett, Ralph Adolphs, Stacy Marsella, Aleix M Martinez, and Seth D Pollak. 2019. Emotional expressions reconsidered: Challenges to inferring emotion from human facial movements. *Psychological science in the public interest* 20, 1 (2019), 1–68.
- [5] Christoph Bartneck, Tony Belpaeme, Friederike Eyssel, Takayuki Kanda, Merel Keijsers, and Selma Šabanović. 2020. *Human-robot interaction: An introduction*. Cambridge University Press.
- [6] RE Belleman. 1961. Adaptive control processes.
- [7] Kush Bhatia, Ashwin Pananjady, Peter Bartlett, Anca Dragan, and Martin J Wainwright. 2020. Preference learning along multiple criteria: A game-theoretic perspective. *Advances in Neural Information Processing Systems* 33 (2020), 7413–7424.
- [8] Gautam Bhattacharya, Koushik Ghosh, and Ananda S Chowdhury. 2017. kNN classification with an outlier informative distance measure. In *International Conference on Pattern Recognition and Machine Intelligence*. Springer, 21–27.
- [9] Erdem Biyik and Dorsa Sadigh. 2018. Batch Active Preference-Based Learning of Reward Functions. In *Proceedings of The 2nd Conference on Robot Learning (Proceedings of Machine Learning Research, Vol. 87)*. Aude Billard, Anca Dragan, Jan Peters, and Jun Morimoto (Eds.). PMLR, 519–528. <https://proceedings.mlr.press/v87/biyik18a.html>
- [10] Jeffrey M Bradshaw, Paul J Feltoovich, and Matthew Johnson. 2017. Human-agent interaction. In *The handbook of human-machine interaction*. CRC Press, 283–300.
- [11] Leo Breiman, Jerome Friedman, Richard Olshen, and Charles Stone. 1984. Classification and regression trees. Wadsworth Int. Group 37, 15 (1984), 237–251.
- [12] Kate Candon, Zoe Hsu, Yoony Kim, Jesse Chen, Nathan Tsoi, and Marynel Vázquez. 2022. Perceptions of the Helpfulness of Unexpected Agent Assistance. In *Proceedings of the 10th International Conference on Human-Agent Interaction (Christchurch, New Zealand) (HAI '22)*. Association for Computing Machinery, New York, NY, USA, 41–50. <https://doi.org/10.1145/3527188.3561915>
- [13] Colleen M Carpinella, Alisa B Wyman, Michael A Perez, and Steven J Stroessner. 2017. The robotic social attributes scale (rosas) development and validation. In *Proceedings of the 2017 ACM/IEEE International Conference on human-robot interaction*. 254–262.
- [14] Cristiano Castelfranchi. 1998. Modelling social action for AI agents. *Artificial intelligence* 103, 1-2 (1998), 157–182.
- [15] Andrea Cuadra, Hansol Lee, Jason Cho, and Wendy Ju. 2021. Look at Me When I Talk to You: A Video Dataset to Enable Voice Assistants to Recognize Errors. *arXiv preprint arXiv:2104.07153* (2021).
- [16] Andrea Cuadra, Shuran Li, Hansol Lee, Jason Cho, and Wendy Ju. 2021. My Bad! Repairing Intelligent Voice Assistant Errors Improves Interaction. *Proceedings of the ACM on Human-Computer Interaction* 5, CSCW1 (2021), 1–24.
- [17] Yuchen Cui, Qiping Zhang, Alessandro Allievi, Peter Stone, Scott Niekum, and W. Bradley Knox. 2020. The EMPATHIC Framework for Task Learning from Implicit Human Feedback. In *Proceedings of the 4th Conference on Robot Learning (CoRL 2020)*. Cambridge MA, USA. <http://www.cs.utexas.edu/users/ai-lab?CORL20-Cui>
- [18] Sanobar Dar and Ulysses Bernardet. 2020. When Agents Become Partners: A Review of the Role the Implicit Plays in the Interaction with Artificial Social Agents. *Multimodal Technologies and Interaction* 4, 4 (2020). <https://doi.org/10.3390/mti4040081>
- [19] Manoranjan Dash and Huan Liu. 1997. Feature selection for classification. *Intelligent data analysis* 1, 1-4 (1997), 131–156.
- [20] Celso M de Melo, Peter J Carnevale, and Jonathan Gratch. 2018. *Peoples Biased Decisions to Trust and Cooperate with Agents that Express Emotions*. Technical Report. UNIVERSITY OF SOUTHERN CALIFORNIA LOS ANGELES.
- [21] Celso M. de Melo, Peter Khooshabeh, Ori Amir, and Jonathan Gratch. 2018. Shaping Cooperation between Humans and Agents with Emotion Expressions and Framing. In *Proceedings of the 17th International Conference on Autonomous Agents and MultiAgent Systems (Stockholm, Sweden) (AAMAS '18)*. International Foundation for Autonomous Agents and Multiagent Systems, Richland, SC, 2224–2226.
- [22] Carmel Domshlak, Eyke Hüllermeier, Souhila Kaci, and Henri Prade. 2011. Preferences in AI: An overview. *Artificial Intelligence* 175, 7-8 (2011), 1037–1052.
- [23] Paolo Dragone, Stefano Teso, and Andrea Passerini. 2018. Constructive preference elicitation. *Frontiers in Robotics and AI* 4 (2018), 71.
- [24] Wen Duan, Naomi Yamashita, Yoshinari Shirai, and Susan R Fussell. 2021. Bridging Fluency Disparity between Native and Nonnative Speakers in Multilingual Multiparty Collaboration Using a Clarification Agent. *Proceedings of the ACM on Human-Computer Interaction* 5, CSCW2 (2021), 1–31.
- [25] David Feil-Seifer and Maja J Mataric. 2005. Defining socially assistive robotics. In *9th International Conference on Rehabilitation Robotics, 2005. ICORR 2005. IEEE*, 465–468.
- [26] Terrence Fong, Illah Nourbakhsh, and Kerstin Dautenhahn. 2003. A survey of socially interactive robots. *Robotics and autonomous systems* 42, 3-4 (2003), 143–166.
- [27] Mark G Frank, Paul Ekman, and Wallace V Friesen. 1993. Behavioral markers and recognizability of the smile of enjoyment. *Journal of personality and social psychology* 64, 1 (1993), 83.
- [28] Helen I. Gallagher, Anthony I Jack, Andreas Roepstorff, and Christopher D Frith. 2002. Imaging the intentional stance in a competitive game. *Neuroimage* 16, 3 (2002), 814–821.
- [29] Manuel Giuliani, Nicole Mirnig, Gerald Stollnberger, Susanne Stadler, Roland Buchner, and Manfred Tscheligi. 2015. Systematic analysis of video data from different human–robot interaction studies: a categorization of social signals during error situations. *Frontiers in psychology* 6 (2015), 931.
- [30] Samuel D Gosling, Peter J Rentfrow, and William B Swann Jr. 2003. A very brief measure of the Big-Five personality domains. *Journal of Research in personality* 37, 6 (2003), 504–528.
- [31] Simon Greipl, Katharina Bernecker, and Manuel Ninaus. 2021. Facial and Bodily Expressions of Emotional Engagement: How Dynamic Measures Reflect the Use of Game Elements and Subjective Experience of Emotions and Effort. *Proceedings of the ACM on Human-Computer Interaction* 5, CHI PLAY (2021), 1–25.
- [32] Jasmin Grosinger, Federico Pecora, and Alessandro Saffiotti. 2019. Robots that maintain equilibrium: Proactivity by reasoning about user intentions and preferences. *Pattern Recognition Letters* 118 (2019), 85–93. <https://doi.org/10.1016/j.patrec.2018.05.014> Cooperative and Social Robots: Understanding Human Activities and Intentions.
- [33] James J Gross, OP John, and J Richards. 1995. *Berkeley expressivity questionnaire*. Edwin Mellen Press Lewiston, NY.
- [34] Balint Gucci, Danesh S Tarapore, William Yeoh, Christopher Amato, and Long Tran-Thanh. 2020. To ask or not to ask: A user annoyance aware preference elicitation framework for social robots. In *2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 7935–7940.
- [35] Luke Guerdan, Alex Raymond, and Hatice Gunes. 2021. Toward affective XAI: facial affect analysis for understanding explainable human-ai interactions. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 3796–3805.
- [36] Haibo He and Edwardo A Garcia. 2009. Learning from imbalanced data. *TKDE* 21, 9 (2009), 1263–1284.
- [37] Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long short-term memory. *Neural computation* 9, 8 (1997), 1735–1780.
- [38] John Houston, Paul Harris, Sandra McIntire, and Dientje Francis. 2002. Revising the competitiveness index using factor analysis. *Psychological Reports* 90, 1 (2002), 31–34.
- [39] Rachael E Jack, Oliver GB Garrod, Hui Yu, Roberto Caldara, and Philippe G Schyns. 2012. Facial expressions of emotion are not culturally universal. *Proceedings of the National Academy of Sciences* 109, 19 (2012), 7241–7244.
- [40] Natasha Jaques, Jennifer McCleary, Jesse Engel, David Ha, Fred Bertsch, Douglas Eck, and Rosalind Picard. 2020. Learning via social awareness: Improving a deep generative sketching model with facial feedback. In *Workshop on Artificial Intelligence in Affective Computing*. PMLR, 1–9.
- [41] Richard Jensen and Qiang Shen. 2009. Are More Features Better? A Response to <emphasis>emphasis-type="italic">Attributes Reduction Using Fuzzy Rough Sets</emphasis>. *IEEE Transactions on Fuzzy Systems* 17, 6 (2009), 1456–1458. <https://doi.org/10.1109/TFUZZ.2009.2026639>
- [42] Hong Jun Jeon, Smitha Milli, and Anca Dragan. 2020. Reward-rational (implicit) choice: A unifying formalism for reward learning. In *Advances in Neural Information Processing Systems*, H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin (Eds.), Vol. 33. Curran Associates, Inc., 4415–4426. <https://proceedings.neurips.cc/paper/2020/file/2f10c1578a0706e06b6d7db6f0b4a6af-Paper.pdf>
- [43] Fernando Koch and Iyad Rahwan. 2004. Classification of agents-based mobile assistants. In *AAMAS workshop on agents for ubiquitous computing*. AAMAS.
- [44] Dimosthenis Kontogiorgos, Minh Tran, Joakim Gustafson, and Mohammad Soleymani. 2021. A Systematic Cross-Corpus Analysis of Human Reactions to Robot Conversational Failures. In *Proceedings of the 2021 International Conference on Multimodal Interaction (Montréal, QC, Canada) (ICMI '21)*. Association for

- Computing Machinery, New York, NY, USA, 112–120. <https://doi.org/10.1145/3462244.3479887>
- [45] Pat Langley et al. 1994. Selection of relevant features in machine learning. In *Proceedings of the AAAI Fall symposium on relevance*, Vol. 184. 245–271.
 - [46] Jamie Large, Graham Stodolski, and Marynel Vázquez. 2020. Studying Human-Agent Interactions in Space Invaders. In *Proceedings of the 8th International Conference on Human-Agent Interaction*. 245–247.
 - [47] J. Lazar, J.H. Feng, and H. Hochheiser. 2017. *Research Methods in Human-Computer Interaction*. Elsevier Science. <https://books.google.com/books?id=hbKxDAQQAQBAJ>
 - [48] Iolanda Leite, André Pereira, Carlos Martinho, Ana Paiva, and Ginevra Castellano. 2009. Towards an empathic chess companion. In *Proceedings of the 8th International Conference on Autonomous agents and Multiagent Systems (AAMAS 2009)*. 33–36.
 - [49] Mark R Lepper, Maria Woolverton, Donna L Mumme, and J Gurtner. 1993. Motivational techniques of expert human tutors: Lessons for the design of computer-based tutors. *Computers as cognitive tools* 1993 (1993), 75–105.
 - [50] Michael Lewis. 1998. Designing for human-agent interaction. *AI Magazine* 19, 2 (1998), 67–67.
 - [51] Guangliang Li, Hamdi Dibeklioglu, Shimon Whiteson, and Hayley Hung. 2020. Facial feedback for reinforcement learning: a case study and offline analysis using the TAMER framework. *Autonomous Agents and Multi-Agent Systems* 34, 1 (2020), 1–29.
 - [52] Jinying Lin, Zhen Ma, Randy Gomez, Keisuke Nakamura, Bo He, and Guangliang Li. 2020. A review on interactive reinforcement learning from human social feedback. *IEEE Access* 8 (2020), 120757–120765.
 - [53] Gale M Lucas, Jonathan Gratch, Aisha King, and Louis-Philippe Morency. 2014. It’s only a computer: Virtual humans increase willingness to disclose. *Computers in Human Behavior* 37 (2014), 94–100.
 - [54] Donald McMillan, Barry Brown, Ikkaku Kawaguchi, Razan Jaber, Jordi Solsona Belenguier, and Hideaki Kuzuoka. 2019. Designing with Gaze: Tama – a Gaze Activated Smart-Speaker. *Proc. ACM Hum.-Comput. Interact.* 3, CSCW, Article 176 (nov 2019), 26 pages. <https://doi.org/10.1145/3359278>
 - [55] Stefan Palan and Christian Schitter. 2018. Prolific. ac—A subject pool for online experiments. *Journal of Behavioral and Experimental Finance* 17 (2018), 22–27.
 - [56] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and E. Duchesnay. 2011. Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research* 12 (2011), 2825–2830.
 - [57] Rosalind W Picard. 2000. *Affective computing*. MIT press.
 - [58] James K Rilling, David A Gutman, Thorsten R Zeh, Giuseppe Pagnoni, Gregory S Berns, and Clinton D Kilts. 2002. A neural basis for social cooperation. *Neuron* 35, 2 (2002), 395–405.
 - [59] Claire Rivoire and Angelica Lim. 2016. The delicate balance of boring and annoying: Learning proactive timing in long-term human robot interaction. (2016).
 - [60] Stephanie Rosenthal, Joydeep Biswas, and Manuela M Veloso. 2010. An effective personal mobile robot agent through symbiotic human-robot interaction.. In *AAMAS*, Vol. 10. 915–922.
 - [61] Aaron Roth, Jonathan Ullman, and Zhiwei Steven Wu. 2016. Watch and learn: Optimizing from revealed preferences feedback. In *Proceedings of the forty-eighth annual ACM symposium on Theory of Computing*. 949–962.
 - [62] Alan G Sanfey, James K Rilling, Jessica A Aronson, Leigh E Nystrom, and Jonathan D Cohen. 2003. The neural basis of economic decision-making in the ultimatum game. *Science* 300, 5626 (2003), 1755–1758.
 - [63] Lindsay Sanneman. 2019. Preference Elicitation and Explanation in Iterative Planning.. In *IJCAI*. 6456–6457.
 - [64] Bill N Schilit and Marvin M Theimer. 1994. Disseminating active map information to mobile hosts. *IEEE network* 8, 5 (1994), 22–32.
 - [65] Albrecht Schmidt. 2000. Implicit human computer interaction through context. *Personal technologies* 4, 2 (2000), 191–199.
 - [66] David C Schrader. 1994. Judgments of acceptable partners for social goals based on perceptions of nonverbal behavior. *Journal of Social Behavior and Personality* 9, 2 (1994), 353.
 - [67] Maia Stiber, Russell Taylor, and Chien-Ming Huang. 2022. Modeling Human Response to Robot Errors for Timely Error Detection. In *2022 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 676–683.
 - [68] Katia Sycara and Gita Sukthankar. 2006. Literature review of teamwork models. *Robotics Institute, Carnegie Mellon University* 31 (2006), 31.
 - [69] Vaibhav V Unhelkar, Shen Li, and Julie A Shah. 2020. Decision-making for bidirectional communication in sequential human-robot collaborative tasks. In *2020 15th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*. IEEE, 329–341.
 - [70] Mascha Van’t Wout, René S Kahn, Alan G Sanfey, and André Aleman. 2006. Affective state and decision-making in the ultimatum game. *Experimental brain research* 169, 4 (2006), 564–568.
 - [71] Hana Vrzakova, Mary Jean Amon, McKenzie Rees, Myrthe Faber, and Sidney D’Mello. 2021. Looking for a Deal? Visual Social Attention during Negotiations via Mixed Media Videoconferencing. *Proc. ACM Hum.-Comput. Interact.* 4, CSCW3, Article 260 (jan 2021), 35 pages. <https://doi.org/10.1145/3434169>
 - [72] Lennart Wachowiak, Peter Tisnikar, Gerard Canal, Andrew Coles, Matteo Leonetti, and Oya Celiktutan. 2022. *Analysing Eye Gaze Patterns during Confusion and Errors in Human-Agent Collaborations*.
 - [73] Ruiqi Wang, Weizheng Wang, and Byung-Cheol Min. 2022. Feedback-efficient Active Preference Learning for Socially Aware Robot Navigation. In *2022 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. 11336–11343. <https://doi.org/10.1109/IROS47612.2022.9981616>
 - [74] Keiichi Yamazaki, Michie Kawashima, Yoshinori Kuno, Naonori Akiya, Matthew Burdelski, Akiko Yamazaki, and Hideaki Kuzuoka. 2007. Prior-to-request and request behaviors within elderly day care: Implications for developing service robots for use in multiparty settings. In *ECSCW 2007*. Springer, 61–78.
 - [75] Yukang Yan, Chun Yu, Wengrui Zheng, Ruining Tang, Xuhai Xu, and Yuanchun Shi. 2020. FrownOnError: Interrupting Responses from Smart Speakers by Facial Expressions. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–14.