

Implicit Communication of Actionable Information in Human-AI teams

Claire Liang, Julia Proft, Erik Andersen, and Ross A. Knepper

Department of Computer Science, Cornell University
cyl48@cornell.edu, {jproft, eland, rak}@cs.cornell.edu

ABSTRACT

Humans expect their collaborators to look beyond the explicit interpretation of their words. *Implicature* is a common form of implicit communication that arises in natural language discourse when an utterance leverages context to imply information beyond what the words literally convey. Whereas computational methods have been proposed for interpreting and using different forms of implicature, its role in human and artificial agent collaboration has not yet been explored in a concrete domain. The results of this paper provide insights to how artificial agents should be structured to facilitate natural and efficient communication of actionable information with humans. We investigated implicature by implementing two strategies for playing *Hanabi*, a cooperative card game that relies heavily on communication of actionable implicit information to achieve a shared goal. In a user study with 904 completed games and 246 completed surveys, human players randomly paired with an implicature AI are 71% more likely to think their partner is human than players paired with a non-implicature AI. These teams demonstrated game performance similar to other state of the art approaches.

CCS CONCEPTS

• **Human-centered computing** → **User studies; Collaborative interaction; Empirical studies in HCI; Empirical studies in collaborative and social computing; Web-based interaction;**

KEYWORDS

Collaboration, Games/Play, Empirical study that tells us about people

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
CHI 2019, May 4–9, 2019, Glasgow, Scotland UK

© 2019 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-5970-2/19/05...\$15.00

<https://doi.org/10.1145/3290605.3300325>

ACM Reference Format:

Claire Liang, Julia Proft, Erik Andersen, and Ross A. Knepper. 2019. Implicit Communication of Actionable Information in Human-AI teams. In *CHI Conference on Human Factors in Computing Systems Proceedings (CHI 2019)*, May 4–9, 2019, Glasgow, Scotland UK. ACM, New York, NY, USA, 13 pages. <https://doi.org/10.1145/3290605.3300325>

1 INTRODUCTION

An important area of human-computer interaction research involves enabling collaborative behavior for artificial agents working in close partnership with humans. Teamwork requires communication for planning and coordination, and the content communicated during this form of coordination is complex and multifaceted. In particular, a significant amount of communication in human groups goes through non-explicit channels [73]. These channels can range from non-explicit verbal statements to other means as subtle as eye gaze or gesture [71].

The impact of implicit communication is ubiquitous. It can be found in applications from robot motion [25] to language input [84]. In many domains involving teamwork, implicit communication is vital for good performance. In order to be effective partners and integrate well with human teams, robots and AIs need to understand and generate this kind of actionable implicit communication that results from actions being situated in context [9, 18, 27].

One area in which implicit communication is particularly important is human-AI interaction in video games. While there has been significant progress in creating AIs that can play games such as chess [16] and Go [78], AIs are still incapable of playing many kinds of games because they cannot effectively deal with social elements [64]. For example, consider the cooperative card game *Hanabi*. In this game, players must cooperate to play a set of cards in a certain order. Each player can see all of the other players' cards, but they cannot see their own. Although players need to inform other players about what cards they should play, they are prohibited from communicating anything other than direct facts such as "this card is red". Therefore, in this highly restricted communication space, players routinely use direct facts to communicate actionable intent *implicitly*, such as saying "this card is red" to convey the action "play this card now" [61]. Capturing implicit information of this form is crucial for designing AIs

that can collaborate with humans in many HCI scenarios such as playing these kinds of games. However, implicit communication is rarely incorporated into AI design, effectively omitting a large component of team game experiences.

Previous work in building AIs that can collaborate with humans, both for *Hanabi* [26, 60] and for other games [77], has generally centered on modeling teammates' mental states and exhaustively searching all potential states. However, these approaches are inefficient. Furthermore, although the artificial agents resulting from these approaches perform well when playing against each other, they reason in a way that most human players do not, resulting in interactions that are insufficiently natural.

Modeling implicit communication is difficult because it is hard to generalize a set of principles that transcends domain and context. For our study, we draw inspiration from the field of pragmatics in linguistics in order to interpret and make use of implicit communication in human conversation. Specifically, this paper takes a step towards this goal by using a type of implicit communication called *conversational implicature* [33]. Our key insight is that in a teamwork setting such as *Hanabi*, implicit communication typically focuses on conveying actionable information.

In this paper, we detail the landscape of communication for teamwork and subdivide the communicated content into different categories of implicit communication. We take a principled approach to studying effective modes of information transfer among teammates by investigating two methods of communicating information. We provide a computational approach to the theory of implicature by using Gricean maxims and describe an AI that we built that uses these strategies to play the cooperative card game *Hanabi*. We conducted a study in simulation comparing the AI that uses implicit communication strategies to a baseline AI that only uses explicit communication; the results of this study demonstrate that the former AI consistently outperforms the latter. Using an online interface of our design, we conducted a user study in which users completed 904 games and 246 follow-up surveys. From this study we found that the AI that used implicit communication performed comparably with other state of the art *Hanabi* AIs in terms of game score, but human teammates were 71% more likely to believe that their partner was human when playing with an AI that used implicit communication.

2 RELATED WORK

In our work, we employ actionable implicit communication as a hint-giving strategy in *Hanabi*. Aspects of implicit communication, though not always actionable, arise in research in many fields. We survey its forms and applications here.

Forms of Implicit Communication

There exists a rich body of work that explores how communication can establish varying forms of shared context or common ground. Ways to reach consensus about context range broadly from demographic information about participating agents [46] to converging on an understanding of game playing strategy for multiplayer online battle arena games [47]. Many of these areas refer to a form of implicit communication, often verbally-restricted or entirely non-verbal, to reach these consensuses.

One form of implicit communication that has been explored in the realm of emotional understanding is *affective grounding*, which tends to focus on converging on a shared knowledge or understanding of mental state [42] in regard to an agent's internal emotional state. Another area in which implicit communication is crucial is *tacit coordination*. Tacit coordination is focused on establishing consensus about expected behavior and responsibilities of others on a team when working to achieve a shared goal [20]. It is cited as an important factor in scenarios from team dynamics for multiplayer game team coordination [88] to effectiveness of crowdsourced teams in achieving tasks [91]. Some means of communication for tacit coordination in these contexts are explicit, but there are others, such as use of an annotation system in distributed multiplayer games, that are highly context-dependent and implicitly communicated [1].

Linguistics literature describing the phenomenon of conversational implicature [11, 35, 37, 38, 51, 67] originated with Grice [33]. Work in natural language processing has sought to build artificial agents able to generate and understand utterances containing implicature. Generalizable techniques are difficult to put into practice in a computational system due to dependency on context representation [72]. More tractable techniques isolate specific and far more restricted types of implicature, such as the technique proposed for scalar implicature [31]. The full complexity of the implicature understanding and generation problems has been explored as well [84] but remains intractable.

Foundations of Implicit Communication

The mechanism of implicit communication, including implicature, leverages context to allow messages to convey more information content than the raw message contains. This is possible because actions (the messages) are interpreted in context, from which added meaning can be derived. In social settings, this mechanism is mediated by an innate human mental capability, called teleological reasoning, which connects goals and actions. People perceive both humans [24] and robots [74] to be goal-oriented. If a person can correctly predict an action that somebody will take to achieve a known goal, we would say that the action is *predictable*; if a person

can infer the goal based on her observation of an action, then we say that the action is *legible* [25]. The fact that humans perform teleological reasoning automatically explains how many *Hanabi* players are able to adopt an implicature-based hinting strategy without explicitly conferring with the other players about strategy.

At an even more basic level, teleological reasoning is a form of *abductive inference* [62], sometimes called “inference to the best explanation”. Just as deduction leads from a general principle ($A \rightarrow B$) and an observation (A) to a consequence (B) and induction leads from a pair of specific observations (A and B) to a prediction about a general principle ($A \rightarrow B$), abduction takes a general principle ($A \rightarrow B$) and an observation (B) and returns an explanation (A). It is an everyday human practice. For example, if Zach notices that his car key is in a different place than he left it, his mind will immediately compare possible explanations: his wife borrowed his car, the cat knocked the key on the floor, or there was an earthquake. Abduction has been studied in AI [13, 15, 30, 36, 41, 50, 58, 63, 65] and employed to solve a variety of real-world problems [36, 57, 66, 81, 82] and to perform plan recognition [5, 21, 85]. Mirroring results in implicature however, it was found that abduction is generally intractable [2, 14, 15, 28, 30] to exactly compute.

Partially Observable Games

There are many strategies for partially observable games [32, 43, 79], but the cooperative nature of *Hanabi* has inspired different kinds of playing strategies. We specifically consider the impact that Gricean conversational implicature can have on *Hanabi*-playing logic [52, 86].

Osawa [60] proposes a series of strategies that consider the state of a teammate’s hand in a two-player game of *Hanabi*. The final strategy, titled the “self-recognition strategy”, outperforms the other strategies by modeling the other player’s knowledge of their hand before providing each hint and considering the inverse when interpreting hints. Osawa implements the self-recognition strategy using an entropy-minimization approach and demonstrates the superior performance of this approach in computer simulations against less advanced models.

Eger et al. [26] uses two of the Gricean maxims (Relation and Manner) in the context of *Hanabi* to construct an AI that employs “intentionality”, using implicature in addressing goal-directedness. Their “intentional” AI also employs an alternative discard policy and other resource management techniques as modifications over Osawa’s prior work. Their AI is geared toward improving the *Hanabi* game experience as well as optimizing the AI for game score performance. Results indicate that for a limited score range, player engagement and enjoyment is correlated with the intentionality level of the AI. However, their AI does not exclusively or

wholly investigate the impact of Grice’s maxims, and their hypotheses are not focused on how conversational implicature affects the extent to which the gameplay of the AI reflects the gameplay of a human. We study the full score range from our user study participants and specifically investigate how the AI is perceived by human players.

We note that there are other rule-based techniques that can create AIs that can successfully demonstrate high scores in *Hanabi*. For example, Bouzy [8] and Cox et al. [23] use ideas from hat guessing games to get near-optimal scores. However, many rule-based approaches apply techniques that would be entirely opaque to human teammates [12, 83].

Implicit Communication in Human-Robot Interaction

Since robots can act physically in the world, they have a greater potential than software agents to perform acts of implicit communication and to best exploit actionable implicit communication from a human. An abstract, high-level framework for robots to automatically understand and generate implicit communicative messages was proposed [48], although as presented it remains intractable to perform this inference for nontrivial problems.

There are many instances of robots performing domain-specific implicit communication. Robots have been demonstrated communicating the weight of objects during lifting [75], the intended goal of a reaching motion [25], how a person can most effectively help the robot [49], intended avoidance maneuvers during socially-competent navigation [56], symbol grounding for language understanding [39], and communicating a self-driving car’s goals to human observers [40].

Researchers have also investigated implicit robot control by programming the robot’s environment [76, 90]. These interfaces qualify as implicit communication in the sense that they leverage context to interpret the meaning of symbols. In order to function as part of an effective team, robots need to infer human intentions and also express the internal state of the robot, for example via affective cues [10, 29].

Conversational Implicature, Games, and HCI

Researchers have examined the role of conversational implicature in human-computer interaction, such as in interpretation of multimodal interactions including eye gaze [68], language [19, 53], and gesture [44]. Cordona-Rivera and Young studied how conversational implicature can provide a framework for game design and communication between the game and the player [17]. We build on this work by applying conversational implicature specifically to human-AI teamwork and transmission of actionable information.

There is some work in HCI that examines board games, focusing on topics such as fostering sustainability [6], collaboration [89], and identification of user interface needs [80]. Rogerson et al. studied *Hanabi* specifically, examining how players cooperate and distribute cognition [71] and how eye tracking can reveal where players are looking and what they are paying attention to [70]. We also study *Hanabi* but focus instead on how to model and improve team dynamics when an AI player is part of the team.

There is also work on using video games as a research platform for large-scale online experiments [3, 4, 54, 55]. We use this experimentation method to study the impact of our implicative AI on a large number of human players.

3 HANABI

The game *Hanabi* is well-suited for our problem because of its collaborative nature and its players' natural tendency to adopt implicative into their playing strategies.

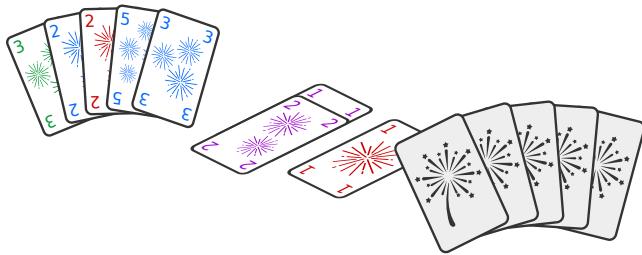


Figure 1: A two-player game of *Hanabi*

Rules

In *Hanabi*, players are a team working towards a shared goal. Although *Hanabi* supports more than two players, we only consider two-player games in this paper. The players alternate turns, creating a symmetric cooperative game. Players start with five cards each and are allowed to look at all players' hands except their own.

The *Hanabi* deck has five colors (suits) and five numbers (ranks). Players also start with three shared fuse tokens that function as lives and eight shared hint tokens (as depicted in Figure 3). For each color there are three cards of rank 1, one of rank 5, and two cards of each rank in between.

On each turn, a player must complete one of the following moves: give a hint, discard a card, or play a card. The goal of the game is to "complete fireworks" by playing cards of all five colors in increasing order from one to five. The game ends if all fireworks are complete, if all fuse tokens are lost, or one round after all cards are drawn from the deck.

Details about each of the three moves are given below.

(1) *Give a hint*: Players can give only two kinds of hints:

- All cards of a color in the teammate's hand (e.g., in Figure 1: "These three cards are blue").
- All cards of a number in the teammate's hand (e.g., in Figure 1: "These two cards are 2s").

The player must point to all applicable cards in the teammate's hand while stating the hint. Giving a hint expends a hint token. Hints can only be given if there are available hint tokens.

- (2) *Discard a card*: The player puts a card into the discard pile face-up, thereby learning the identity of the card. The player then draws a new card from the deck. If any hint tokens are currently spent, one is earned back.
- (3) *Play a card*: The player picks a card from her hand to play and sets it on the table face-up. The card is either
 - next in sequence (a legal play) and is added to its respective stack, or
 - not next in sequence and is added to the discard pile. One fuse token is lost.

The player does not need to specify which sequence she thinks the card belongs to before playing the card. Regardless of whether the play was successful or not, the player draws a new card from the deck. Successfully playing a card of rank 5 also restores a hint token.

At the end of the game, the score is the number of cards that were successfully played. It can range from 0 to 25.

Gameplay

In *Hanabi*, the most significant contextual facts are the identities of currently-playable cards and the cards already played. In addition, it is common knowledge that the most direct way of sharing actionable information with another player (that is, telling him which card to play) is forbidden by the rules. By acknowledging these common-knowledge facts, the team is often able to implicitly communicate the same actionable information by a legal hint.

For example, in Figure 1, suppose that Bernard is the one with the visible hand and that Annie is providing him a hint. Suppose Annie says "You have one red card" to Bernard and indicates a card in his hand about which he has received no previous hints. Even an intermediate *Hanabi* player immediately recognizes this hint as "play the red card". Bernard quickly reasons about this hint. He realizes that providing this hint costs resources (a game turn and a hint token) and that Annie will want to pack immediately usable information into it. In addition, given the current game state, it is probabilistically unlikely for Annie to provide a hint about only one card unless she intends for it to be playable. This type of thought process is intuitive, and new players pick up on it as they play. However, in order to create an AI that is capable of this form of implicit communication, we need to formulate a methodology that is suited for a computational

approach. We can use an established theory from the field of pragmatics, the linguistics subfield that studies language in context, and adapt it to develop a policy for this form of implicit communication.

4 CONVERSATIONAL IMPLICATURE

The term *implicature*, coined by Grice [33], is used to describe information that is not explicitly stated but is intentionally conveyed. Grice focused primarily on implicature that occurs in conversation due to what he calls the *cooperative principle*, meaning that speakers are expected to contribute what is required by the accepted purpose of the conversation [87]. He gives four maxims that constitute the cooperative principle and describe how to conduct cooperative speech [34]:

- Quality: only contribute information that is true.
- Quantity: provide all necessary information, but not more.
- Relation: make your contribution relevant.
- Manner: avoid ambiguity; be clear, orderly, and brief.

Implicature arises when any of the maxims are violated. Consider the following example from Recanati [69]:

- Annie: “Can you cook?”
- Bernard: “I am French.”

In order to make sense of Bernard’s response, Annie must apply the following deductive inference steps:

- (1) *Contextual premise*: Bernard is able to answer the question of whether he can cook.
- (2) *Contextual premise*: It is common knowledge that the French are known for their ability to cook.
- (3) Assume Bernard adheres to the cooperative principle and the four maxims.
- (4) By (1), Bernard can completely resolve Annie’s question, and by (3), he will.
- (5) Only the propositions that Bernard can or cannot cook can fully resolve the question.
- (6) By stating a fact seemingly irrelevant to the propositions of (5), Bernard flouts the maxim of Relation.
- (7) Thus Annie must search over a plausible set of facts in the relevant common ground to find fact (2) and conclude that Bernard is implicating that he is able to cook.

There is also a well-known cancellability test to gauge whether an instance of implicature is conversational implicature [7]. The reasoning in Lines (4)–(7) depends on the concepts of shared trust, consensus in a team, and receptivity of teammates. These dependencies are described in greater detail by Knepper et al. [48].

We can then apply this set of rules to *Hanabi*.

5 AI DESIGN

The core strategy in *Hanabi* occurs during the giving and interpretation of hints. A key distinction in the strategic use of hints involves the issue of who decides what cards to play. If the team is viewed as a set of individual decision-makers, then each player is in charge of her own hand and actions. In this view, the best hinting strategy is to maximize information about the teammate’s hand so that she can make the best decisions. We call the AI that uses this hint strategy the entropy AI, which serves as our baseline.

A contrasting hinting strategy stems from the perspective that good teamwork requires working together. Since each player has complete knowledge of her teammate’s cards, it is logical that each player should decide what card her teammate should play next. In this case, the hints serve to direct the teammate’s attention to which (singleton) card should be played next via implicature. We call the AI that uses this hint strategy the implicature AI.

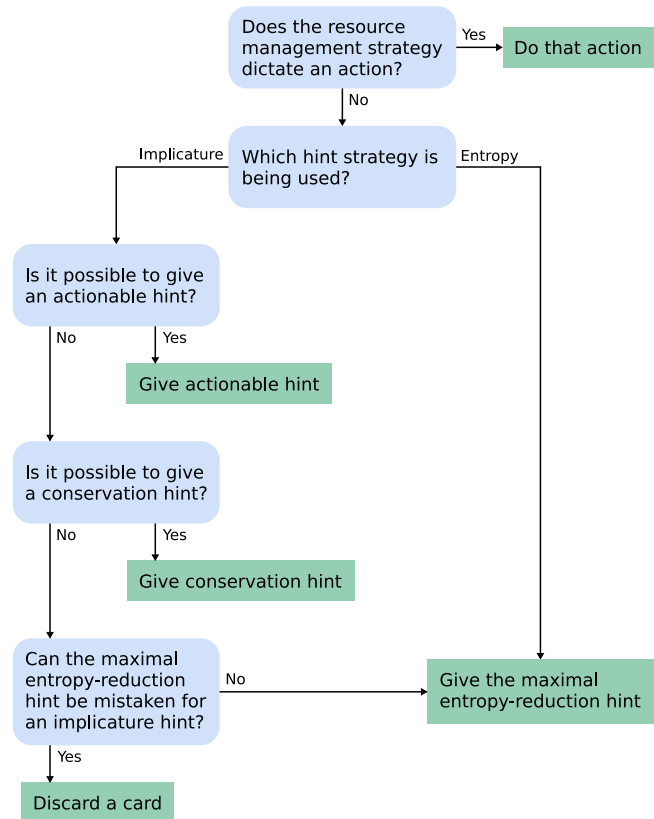


Figure 2: Decision tree illustrating the logic of the different hint strategies.

The overall AI structure is shown in Figure 2. The AIs share significant gameplay infrastructure, termed “resource management strategy” in the figure. The shared strategy is

as follows. First, if the AI knows that it can play a card, it does so. Second, there is logic to special-case the hinting and playing of 1s, since it is desirable to give a maximal entropy-reduction hint about 1s even in the implicature condition. Third, the AI considers whether it should discard a card. If it knows for certain that one of its cards can be safely discarded (for example, if a blue 2 has already been played successfully, and it knows that it is holding the second blue 2), it will discard that card. Otherwise, if it is not certain whether any cards in its hand can be safely discarded but there is only one hint token left, it tries to discard a card that it suspects can be safely discarded. Lastly, if there are no hint tokens available, it will discard the card that it knows least about. From there, the AI takes a different course of action depending on which hint strategy the AI implements.

Implicature AI

Broadly, the logic of the implicature hint-giving strategy proceeds as described in Section 4. According to that section, the actor triggers an implicature interpretation by flouting one of the Gricean maxims. Since the observer assumes that the actor is being cooperative, she infers that there must exist some hypothetical actionable information that is unknown to her but would be known to the actor. She then searches over possible messages containing actionable information that the actor could be implicating. With a good choice of actionable information, a maxim violation can be interpreted as a cooperative behavior instead.

The implicature AI's connection to the Gricean maxims is based on several observations. First, giving hints uses up hint tokens, making it a costly move not to be made frivolously. Second, circumstances in *Hanabi* do not always permit implicature hints: either there may be no card with actionable information or there may be no way to isolate the card by color or number. The hinter may need to fall back on the explicit hint strategy, which may not be worth the hint token. Players need to detect if a given hint is intended to be interpreted explicitly or implicitly. A Gricean maxim applies here based on a statistical argument. Under an explicit entropy-reduction strategy, the best hint on average reveals the color or number of about 2.5 cards, and about 99.85% of hands permit a hint about two or more cards. An implicature hinting strategy should be able to give more immediately useful and therefore more efficient hints by communicating actionable information. Thus, a hint about a single card flouts the maxim of Quantity, so the hint implicates actionable information rather than merely conveying identity information.

The implicature AI does not perform Gricean reasoning directly. Instead, we built a decision tree that encodes the application of the Gricean maxims to giving and receiving hints (see Figure 2). The first two of the three decisions in the implicature hint strategy involve Gricean reasoning, and the

third decision ensures that unintended implicit meaning is not conveyed. Details of the three decisions are given below.

- (1) The first Gricean decision addresses whether the AI can give an actionable hint that would cause its teammate to successfully play a card. Qualifying conditions are that the teammate must have a playable card that she does not know about and that there must be a hint that can be given about that card alone. The AI will give that hint if these conditions are met.
- (2) The second Gricean decision involves conservation hints, which are given to prevent a player discarding a valuable card rather than to play a card. In our implementation, only 5s are considered for conservation hints since they are each unique in the deck. If the teammate is holding a 5 that she is unaware of, then a conservation hint is given.
- (3) Finally, the implicature AI takes care to account for the maxims by only providing a hint about a singleton card if that card is actionable or unambiguously conservable. As a result, the AI will only give a maximal entropy-reduction hint if the hint concerns more than one card, thereby avoiding hints that could be misinterpreted as a suggestion to take an action with a card.

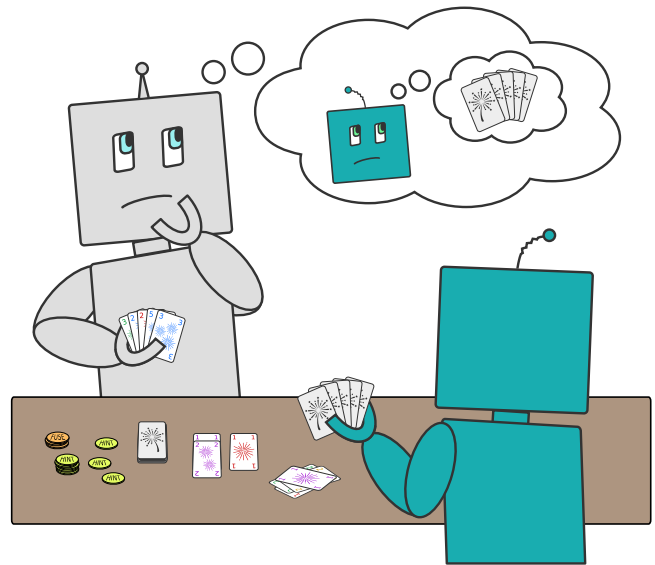


Figure 3: When generating a hint, the AI simulates how the teammate will interpret the hint based on the shared context and what it knows about the teammate's knowledge. For example, the light gray robot, using the implicature AI, might think: "My teammate's leftmost card is a purple three. My teammate knows nothing about it. We both know that the purple three is next in sequence. Since my teammate has only one three, if I hint that it is a three, he will infer that it must be the purple three and will therefore play it."

In order to interpret hints, the AI assumes its teammate uses the same process of reasoning and therefore the same hint strategy. The AI can then use the decision tree as a way to simulate its teammate’s mental model and make sense of the hints it receives. Figure 3 depicts the way the AI simulates its teammate’s thought process to determine whether there is implicature behind the hint [84]. The AI is unable to see its own hand and therefore relies entirely on the hints its teammate provides to narrow down the identity of its cards.

The AI knows for each card in its hand which colors and which numbers are possible given the current state of the game. When the AI receives a hint, it considers each possible hand and the flow of the decision tree in Figure 2 to determine what kind of hint it would give if its teammate was the one with that possible hand. The AI can then use this process to determine the likely identities of the cards in its hand by checking if the hint it would have given aligns with the hint it actually received from its teammate. For example, if the AI receives the hint “you have one 2”, then only the possibility of the hinted card being actionable aligns with this hint. All possible hands that would result in this hint being generated require that the card be playable. Therefore the AI interprets the hint it received as “play this card.”

Entropy AI

As a baseline comparison to the implicature AI, we created an entropy-minimization AI that reflects the hinting strategy of prior work without introducing confounding factors.

An obvious strategy for *Hanabi* is to view one’s partner’s ignorance about the identity of his cards as the primary obstacle. The more a player knows about his own hand, the more likely he can make an informed decision about the appropriate move to take with it.

We can think of this strategy in terms of entropy minimization. The less one knows about a card, the higher the entropy, so attempting to minimize entropy in the context of *Hanabi* is equivalent to trying to inform one’s partner about as many cards as possible per hint.

For example, suppose in Figure 1 that Bernard is the one with the visible hand and Annie has to provide information:

- (1) Annie wants Bernard to be able to make the best move he can from his own knowledge.
- (2) Bernard has no knowledge about his current hand.
- (3) Annie’s goal is to maximize the difference in knowledge between his current hand and his hand at the next turn. That is, she seeks to minimize how much he does not know as of the next turn.
- (4) The maximal entropy-reduction hint Annie can give Bernard is about blue cards (since he has three).

We observe that the entropy strategy fundamentally involves the player with more knowledge (the hint-giver) entrusting her partner to make a good decision about which card to play. The bottleneck is the lack of information about the identities of the hint recipient’s cards.

A shortcoming of this hint strategy is that hints intended only to minimize entropy are unnatural for human *Hanabi* players as they are not actionable information. If Bernard receives the hint that he has three blue cards, he will not know what to do with those blue cards and therefore cannot derive his next move from this hint alone.

6 EVALUATION HYPOTHESES

The goal of our study is to address two hypotheses about the use of implicature in *Hanabi*.

- H1. Teams playing with the implicature AI achieve overall better scores than teams playing with the entropy AI.
- H2. Humans will perceive the implicature AI to be more human-like than the entropy AI.

7 EVALUATION WITH AI-AI GAMES

This section describes a simulation-only study designed to explore hypothesis H1.

Simulation Setup and Execution

We evaluated each hint-giving strategy described above by teaming each AI with a copy of itself. We ran 1000 trial games with the implicature AI-AI team and 1000 trial games with the entropy AI-AI team, using final score as our metric for collaborative effectiveness. H1 predicts a higher average final score for the implicature AI than for the entropy AI.

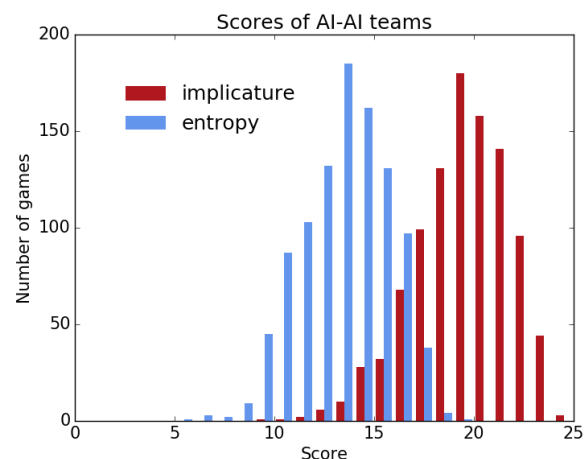


Figure 4: Final scores of games played by an AI-AI team employing either the implicature-based strategy (red) or the entropy-minimization strategy (blue).

Results

Our AI-AI results (Figure 4) show that our implicature-based AI achieves scores that are on average 46% higher than the AI employing entropy minimization (18.9 points on average for implicature games and 13.0 points on average for entropy games). A Student's t -test with two-sample unequal variance is significant ($t = 57.967, p < 0.001$). The Cohen's D value is 2.592381.

Discussion of Simulation Results

Both AIs are implemented statically, so they do not adapt to their teammate's playing style. Each AI played with a teammate that had the exact same hinting strategy, and it is likely that this hard-coded consensus is a big factor in good score performance. However, in reality, human players are not static: they adapt to their teammate's playing strategy. In order to truly investigate the impact of using implicature in our AI design, it is necessary to explore whether non-static agents demonstrate a difference in performance similar to that of our statically-implemented AIs.

8 EVALUATION WITH HUMAN PLAYERS

In this section, we present a web-based human user study designed to explore both hypotheses. We study how humans perform in play with and perceive the AIs.

A dilemma with humans *Hanabi* players is that they often communicate through side-channels during gameplay, such as by conveying information through facial expression. In order to focus strictly on implicit communication conveyed through hints, we created an online environment for two-player games that inherently removed these other forms of implicit communication from the game experience.

Interface

The interface presented in Figure 5 is meant to be as close to natural in-person gameplay as permitted by the *Hanabi* rules. In addition, we intentionally designed the interface to isolate communication of informational content, removing emotion from the game entirely. Due to its online nature, we added additional features to compensate for some of the advantages in real-life games.

When the game is administered, the user is directed to a starting lobby. After deciding to start the game, the user is automatically paired up with a "teammate" who is selected randomly from one of our two AIs. We divulge nothing about the teammate's identity (whether it is a human or an AI) nor do we prompt them to consider this when playing the game; the user can only observe their teammate's selected moves as the game unfolds. Whether the user or the teammate goes first is randomly chosen. After the user is paired with their

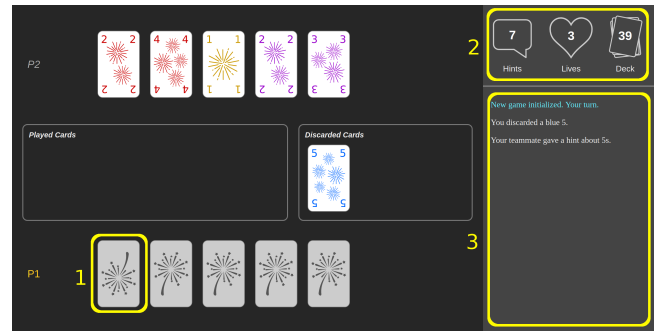


Figure 5: Online *Hanabi* playing interface. An example of a rotated card is shown in (1). The remaining hint tokens, fuse tokens, and cards in the deck are given by (2). The history of all moves made by both players throughout the game is recorded in (3).

teammate and the game begins, they are presented with the interface shown in Figure 5.

When it is the player's turn, they can either select a card to play or discard, or they can provide a hint to their teammate. When receiving a hint from a teammate, the cards that are being hinted about are highlighted. All plays, discards, and hints are recorded in the log on the interface. The interface also allows players to rotate cards in their hand and shuffle them around in case they want to use orientation or ordering of their cards to recall information. However, the other player cannot see how the cards have been moved, so this feature is solely used for keeping track of one's own knowledge.

Aside from playing cards, discarding cards, and giving hints, players are not provided any means communication with their teammate. As a result, there is no side-channel communication that could introduce confounding factors.

User Study Setup and Execution

We recruited participants for our user study through online forums, namely from Reddit (r/boardgames and r/cardgames), and linked them to our online *Hanabi* implementation. Our user study consisted of 904 completed games of *Hanabi* from online users, 463 of whom played with our implicature-based AI and 441 of whom played with our entropy-based AI. All users were prompted to continue to our survey after completing their first game, and we collected 246 completed surveys (127 of these survey-takers played with the implicature AI and 119 with the entropy AI).

We considered several metrics, the most important being game score. However, we also considered the ratio of game score to number of hints given (to measure the efficiency of each hint) as well as the number of fuse tokens used and the number of discarded cards that were playable at the time they were discarded. H1 predicts the overall score for games

played with the implicature AI to be higher. It also predicts a higher score/hints ratio, fewer fuse tokens used, and fewer discarded playable cards.

Our survey asked questions about users' previous experience with *Hanabi*, their opinions on the interface, and their experience overall. Users were asked to rank on a five-point Likert scale if their perception of their team's performance matched up with their overall game score. We also asked whether the player understood their partner's moves and if they thought that their partner understood their moves. H1 predicts that teammates should understand one another better with the implicature AI than with the entropy AI. Lastly, we asked players if they thought their teammate was a human player or a computer player. H2 predicts that the implicature AI should more frequently pass for a human than the entropy baseline.

Results

We found no significant differences in any of the quantitative metrics between users who played with the implicature AI and users who played with the entropy AI; these metrics were score ($t = 0.327, p = 0.74$), ratio of game score to number of hints given ($t = 1.396, p = 0.16$), total fuse tokens used ($t = 1.530, p = 0.13$), and number of playable discards ($t = 0.391, p = 0.70$). Figure 6 shows the final scores for all games that were completed. The ANCOVAs using *Hanabi* experience as a covariate found a p value of 0.442 for scores between implicature and entropy teams. Therefore, there is still no significant effect of hinting strategy on score even after accounting for *Hanabi* experience.

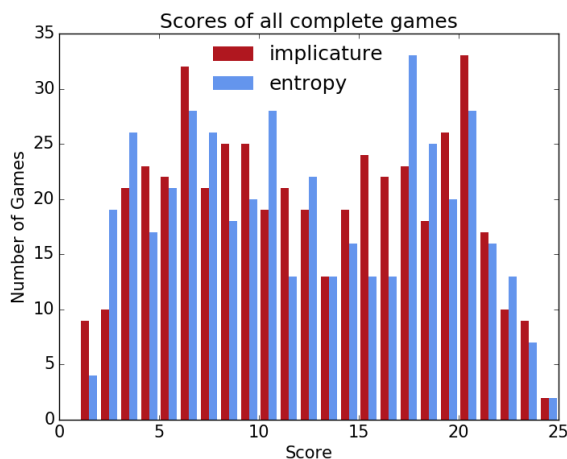


Figure 6: Final scores of games played by a human-AI team in which the AI employed either the implicature-based strategy (red) or the entropy-minimization strategy (blue).

When divided into groups of implicature-AI players and entropy-AI players, a chi-squared test found the group that

played with our implicature AI was significantly more likely to consider our AI to be a human teammate when presented with the options of “human” and “computer” ($X^2 = 7.642, df = 1, p = .0057$). Out of participants who played with our implicature AI, 41% responded that they thought their teammate was human, while only 24% of participants paired with the entropy AI reported the same (see Figure 7).



Figure 7: Comparison of responses to “Do you think your teammate was a human or computer player?” between implicature-teammate and entropy-teammate subgroups.

User Study Discussion

Failure cases for the entropy AI. Some comments from our *Hanabi* players who were assigned to the entropy AI specifically referenced the lack of necessity for complete information due to implicature. One user said, “*Hints are normally given for a reason - You don’t need complete information to play a card.*” A different user cited an instance that referred to the urgency of immediately actionable cards: “*The computer needs to let the player know when they have playable cards by giving timely hints.*” A user even explained the notion of flouting maxims of Relevance and Quantity in their frustration with the AI: “*If someone gives an unclear hint about only one card, that’s an indication to PLAY that card.*”

Failure cases for the implicature AI. Some players paired with the implicature AI noted their confusion about our AI’s playing assumptions. In one failure case, the user hinted to the AI that they had a 1, which the AI took to mean that they should play the card when the user actually intended the opposite (because a card of rank 1 has a low probability of being playable when four suits are already on the table). They explained in the survey, “*for the 1 cards, if a lot are on the table, ideally the card should be discarded not played.*” A different user had the same concern: “*It seemed like any hints I gave were interpreted as ‘play this card.’ I tried to give a hint about 1s when we had four 1s out (so my partner would discard the three 1s in their hand), but my partner ended up just playing 1s instead.*” These instances are interesting because there also exist cases where the hint “you have one 1” would seem like an obvious hint to play the card when there is only one unplayed suit remaining. The issue here is a disagreement of the context that the human player and AI consider when interpreting a hint using implicature.

Implications. The results from this paper demonstrate the feasibility of implicature computationally in real collaborative environments. Although implicature is not sufficient

to capture realistic and natural collaboration, our presented results both show the positive impact that implicature has on a human teammate’s perception of the interaction and point toward the next factors to consider to create an AI that can collaborate effectively with humans.

Both AIs were implemented statically; they did not adapt to their teammate’s preferences and playing style. This choice was intentional to ensure that we could isolate the impact of implicature in this domain without confounding factors. The failure cases of the implicature AI could be addressed by reevaluating the assumed game strategy given the teammate’s moves. For example, one incorrect play of a card that sacrifices a fuse token is enough to realize that there is a discrepancy in how teammates may prioritize their cards.

The next frontier in advancing *Hanabi* game AIs is addressing adaptability. Nikolaidis et al. [59] create an adaptive AI that is able to update its strategy when interpreting and using implicature, which would likely outperform our statically-implemented AIs. Implicature and adaptivity particularly suit combined use since adaptivity can update strategy for effective implicature use. Future work will need to employ adaptivity to progress toward a truly collaborative AI.

Prior work showed some evidence that intentionality may make players enjoy the game experience more [26]. We examine the impact of Grice’s maxims further and specifically determine that implicature leads players to think their AI partner is a human. This result has broad implications for human-AI teamwork design and game AI design because it demonstrates that these maxims, and perhaps the use of pragmatics in general, are a key part of what human players use to evaluate whether their teammates think like humans.

The use of implicature extends to many applications beyond game playing, such as a robot that assembles IKEA furniture alongside a human teammate. Using implicature, a robot working with a human who needs to hammer in a nail can correctly interpret “Hand me *that* red thing” as “Hand me the red hammer closest to you because it is relevant to my current task”. The same logic can extend further to incorporate nonverbal means of communication as well, such as gesture or eye gaze.

9 CONCLUSION

Our work investigates the impact of Gricean conversational implicature from pragmatics in linguistics as instantiated in an AI to play the cooperative game *Hanabi*. We specifically explore (1) whether teams that leverage implicature demonstrate overall better scores than teams that do not, and (2) whether human teammates perceive the implicature-based AI to be more human-like than the baseline AI. We design and implement the proposed implicature AI and the baseline AI to evaluate performance in terms of these two questions. Our simulation results show that implicature-based AI dyads

outperform entropy-based AI dyads in score by 46%, and the results from our online user study indicate that humans are 71% more likely to believe their teammate to be human if the teammate employs our implicature hint-giving strategy.

Our user study contributes human-provided evidence that *Hanabi* players are more likely to find their teammates to be more human-like when implicature is used. In addition, our simulation study suggests that implicature has the potential to enhance team performance, although the result was not replicated in the user study. While implicature on its own is not sufficient for an entirely convincing AI teammate, it is a crucial component to consider when building the collaborative AIs of the future.

Although our AIs and studies focused on *Hanabi* gameplay, the principles of generating and understanding actionable information conveyed via implicit communication have much broader implications. Theoretical frameworks have been proposed to describe the generic mechanism of implicit communication [48] and the mental model for an agent performing a joint activity as part of a team [22]. When combined, these frameworks suggest possible uses for actionable implicit communication at two levels: (1) at the low level to mediate the mechanics of team interactions, and (2) at the high level to furnish pragmatic competence in order to facilitate achievement of shared goals.

Teams need to communicate in order to establish and maintain common ground for the purpose of agreeing on joint intentions and shared goals. At this low level, actionable implicit communication can mediate a wide variety of teamwork behaviors. A *joint intention* within a team is a commitment both to coordinate performance of some action and to inform one another about its progress. The mechanism of actionable implicit communication naturally arises to coordinate joint intentions with teammates since it efficiently encodes that information atop functional behaviors that are part of the joint activity. For example, when a robot is helping a human assemble Ikea furniture [49], the human might pick up an Allen wrench. This action signals to the robot the initiation of a private intention to install fasteners. The robot responds by picking up suitable screws and readying them to be installed, thus signalling to the human the consensus of a joint intention to collaboratively install the screws. In this manner, a considerable amount of coordinating communication during manual teamwork is actionable, implicit, nonverbal, and extremely efficient.

Teams can also leverage actionable implicit communication deliberately to convey information about the task being performed. Effective communication with human collaborators requires *pragmatic competence*, the ability to understand the various ways in which an action will be interpreted by an

observer, and utilize this to achieve desired goals [45]. A robot language tutor could be programmed to select maximally-didactic actions in response to errors made by the human language learner. For example, if a human who is learning Spanish mixes up verbs, she might say *toma el bloque* (take the block) when she means *dame el bloque* (give me the block). The learner expects the robot to pick up the block and hand it over. If the robot instead keeps the block for itself, then the surprise reveals to the learner a contradiction between what she said and what she meant. By forcing the learner to confront the contradiction, the robot causes her to search for an explanation in her word choice.

These examples illustrate the enormous potential for actionable implicit communication, in areas as diverse as robotics, intelligent agents, game AIs, and personal home assistants. The impact of conversational implicature as addressed in this work points to the necessity to explore the next frontier: *pragmatic competence for joint intention*. Tackling this set of problems will be crucial for developing socially-fluent artificial agents.

ACKNOWLEDGMENT

This material is based upon research supported by the Office of Naval Research under Award Number N00014-16-1-2080. We are grateful for this support.

REFERENCES

- [1] Sultan A Alharthi, Ruth C Torres, Ahmed S Khalaf, Zachary O Touns, Igor Dolgov, and Lennart E Nacke. 2018. Investigating the impact of annotation interfaces on player performance in distributed multiplayer games. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. ACM, 314.
- [2] Dean Allemang, Michael C Tanner, Tom Bylander, and John Josephson. 1987. Computational complexity of hypothesis assembly. In *Proceedings of the 10th International Joint Conference on Artificial Intelligence*, Vol. 2. Morgan Kaufmann Publishers Inc., 1112–1117.
- [3] Erik Andersen, Yun-En Liu, Rich Snider, Roy Szeto, and Zoran Popović. 2011. Placing a value on aesthetics in online casual games. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, 1275–1278.
- [4] Erik Andersen, Eleanor O'Rourke, Yun-En Liu, Rich Snider, Jeff Lowdermilk, David Truong, Seth Cooper, and Zoran Popovic. 2012. The impact of tutorials on games of varying complexity. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, 59–68.
- [5] Douglas E Appelt and Martha E Pollack. 1992. Weighted abduction for plan ascription. *User Modeling and User-Adapted Interaction* 2, 1-2 (1992), 1–25.
- [6] Amartya Banerjee, Michael S Horn, and Pryce Davis. 2016. Invasion of the energy monsters: A family board game about energy consumption. In *Proceedings of the 2016 CHI Conference Extended Abstracts on Human Factors in Computing Systems*. ACM, 1828–1834.
- [7] Michael Blome-Tillmann. 2008. Conversational implicature and the cancellability test. *Analysis* 68, 2 (2008), 156–160.
- [8] Bruno Bouzy. 2017. Playing Hanabi near-optimally. In *Advances in Computer Games*. Springer, 51–62.
- [9] Cynthia Breazeal, Cory D Kidd, Andrea Lockerd Thomaz, Guy Hoffman, and Matt Berlin. 2005. Effects of nonverbal communication on efficiency and robustness in human-robot teamwork. In *Intelligent Robots and Systems (IROS), 2005 IEEE/RSJ International Conference on*. IEEE, 708–713.
- [10] Cynthia Breazeal and Brian Scassellati. 1999. How to build robots that make friends and influence people. In *Intelligent Robots and Systems, 1999. IROS'99. Proceedings. 1999 IEEE/RSJ International Conference on*, Vol. 2. IEEE, 858–863.
- [11] Richard Breheny, Napoleon Katsos, and John Williams. 2006. Are generalised scalar implicatures generated by default? An on-line investigation into the role of context in generating pragmatic inferences. *Cognition* 100, 3 (2006), 434–463.
- [12] Cameron B Browne, Edward Powley, Daniel Whitehouse, Simon M Lucas, Peter I Cowling, Philipp Rohlfshagen, Stephen Tavener, Diego Perez, Spyridon Samothrakis, and Simon Colton. 2012. A survey of Monte Carlo tree search methods. *IEEE Transactions on Computational Intelligence and AI in Games* 4, 1 (2012), 1–43.
- [13] Bruce Buchanan, Georgia Sutherland, and Edward A Feigenbaum. 1969. Heuristic DENDRAL: A program for generating explanatory hypotheses in organic chemistry. In *Machine Intelligence*, Bernard Meltzer and Donald Michie (Eds.). Vol. 4. Edinburgh University Press.
- [14] Tom Bylander. 1991. The monotonic abduction problem: A functional characterization on the edge of tractability. In *Proceedings of the Second International Conference on Principles of Knowledge Representation and Reasoning*. Morgan Kaufmann Publishers Inc., 70–77.
- [15] Tom Bylander, Dean Allemang, Michael C Tanner, and John R Josephson. 1991. The computational complexity of abduction. *Artificial Intelligence* 49, 1-3 (1991), 25–60.
- [16] Murray Campbell, A Joseph Hoane Jr, and Feng-hsiung Hsu. 2002. Deep Blue. *Artificial Intelligence* 134, 1-2 (2002), 57–83.
- [17] Rogelio Enrique Cardona-Rivera and Robert Michael Young. 2014. Games as conversation. In *Tenth Artificial Intelligence and Interactive Digital Entertainment Conference*.
- [18] Cristiano Castelfranchi, Giovanni Pezzulo, and Luca Tummolini. 2012. Behavioral implicit communication (BIC): Communicating with smart environments. *Innovative Applications of Ambient Intelligence: Advances in Smart Systems: Advances in Smart Systems* (2012), 1.
- [19] Joyce Yue Chai, Zahar Prasov, and Shaolin Qu. 2006. Cognitive principles in robust multimodal interpretation. *Journal of Artificial Intelligence Research* 27 (2006), 55–83.
- [20] Huo-Tsan Chang, Cheng-Chen Lin, Cheng-Hung Chen, and Yeong-Ho Ho. 2017. Explicit and implicit team coordination: Development of a multidimensional scale. *Social Behavior and Personality: an international journal* 45, 6 (2017), 915–929.
- [21] Eugene Charniak and Robert Goldman. 1991. A probabilistic model of plan recognition. In *Proceedings of the Ninth National Conference on Artificial Intelligence*, Vol. 1. AAAI Press, 160–165.
- [22] Philip R Cohen and Hector J Levesque. 1991. Teamwork. *Nous* 25, 4 (1991), 487–512.
- [23] Christopher Cox, Jessica de Silva, Philip Deorsey, Franklin H. J. Kenter, Troy Retter, and Josh Tobin. 2015. How to make the perfect fireworks display: Two strategies for Hanabi. *Mathematics Magazine* 88 (2015), 323–336.
- [24] G. Csibra and G. Gergely. 2007. 'Obsessed with goals': Functions and mechanisms of teleological interpretation of actions in humans. *Acta Psychologica* 124, 1 (2007), 60–78.
- [25] Anca D Dragan, Kenton CT Lee, and Siddhartha S Srinivasa. 2013. Legibility and predictability of robot motion. In *Human-Robot Interaction (HRI), 2013 8th ACM/IEEE International Conference on*. IEEE, 301–308.
- [26] Markus Eger, Chris Martens, and Marcela Alfaro Córdoba. 2017. An intentional AI for Hanabi. In *CIG 2017*.

- [27] Elliot E. Entin and Daniel Serfaty. 1999. Adaptive Team Coordination. *Human Factors* 41, 2 (1999), 312–325.
- [28] Olivier Fischer, Ashok Goel, John R Svirbely, and Jack W Smith. 1991. The role of essential explanation in abduction. *Artificial Intelligence in Medicine* 3, 4 (1991), 181–191.
- [29] Rachel Gockley, Reid Simmons, and Jodi Forlizzi. 2006. Modeling affect in socially interactive robots. In *Robot and Human Interactive Communication, 2006. ROMAN 2006. The 15th IEEE International Symposium on*. IEEE, 558–563.
- [30] Ashok K Goel, John R Josephson, Olivier Fischer, and P Sadayappan. 1995. Practical abduction: Characterization, decomposition and concurrency. *Journal of Experimental & Theoretical Artificial Intelligence* 7, 4 (1995), 429–450.
- [31] Noah D Goodman and Andreas Stuhlmüller. 2013. Knowledge and implicature: Modeling language understanding as social cognition. *Topics in Cognitive Science* 5, 1 (2013), 173–184.
- [32] Thore Graepel, Ralf Herbrich, and Julian Gold. 2004. Learning to fight. In *Proceedings of the International Conference on Computer Games: Artificial Intelligence, Design and Education*. 193–200.
- [33] Herbert P Grice. 1975. Logic and conversation. *Syntax and Semantics* (1975), 41–58.
- [34] H Paul Grice. 1978. Further notes on logic and conversation. 1978 1 (1978), 13–128.
- [35] Julia Linn Bell Hirschberg. 1985. *A theory of scalar implicature*. University of Pennsylvania.
- [36] Jerry R Hobbs, Mark E Stickel, Douglas E Appelt, and Paul Martin. 1993. Interpretation as abduction. *Artificial Intelligence* 63, 1-2 (1993), 69–142.
- [37] Laurence Horn. 1984. Toward a new taxonomy for pragmatic inference: Q-based and R-based implicature. *Meaning, Form, and Use in Context: Linguistic Applications* (1984), 11–42.
- [38] Laurence R. Horn. 2004. Implicature. In *The Handbook of Pragmatics*, Laurence R. Horn and Gregory Ward (Eds.). Blackwell, Oxford, 3–28.
- [39] Julian Hough and David Schlangen. 2017. A model of continuous intention grounding for HRI. (2017).
- [40] Sandy H Huang, David Held, Pieter Abbeel, and Anca D Dragan. 2017. Enabling robots to communicate their objectives. *Autonomous Robots* (2017), 1–18.
- [41] John R Josephson, B Chandrasekaran, Jack W Smith, and Michael C Tanner. 1987. A mechanism for forming composite explanatory hypotheses. *IEEE Transactions on Systems, Man, and Cybernetics* 17, 3 (1987), 445–454.
- [42] Malte F Jung. 2017. Affective grounding in human-robot interaction. In *Proceedings of the 2017 ACM/IEEE International Conference on Human-Robot Interaction*. ACM, 263–273.
- [43] Leslie Pack Kaelbling, Michael L Littman, and Anthony R Cassandra. 1998. Planning and acting in partially observable stochastic domains. *Artificial intelligence* 101, 1 (1998), 99–134.
- [44] Edward C Kaiser. 2013. Systems and methods for implicitly interpreting semantically redundant communication modes. US Patent 8,457,959.
- [45] Gabriele Kasper. 1997. Can pragmatic competence be taught? <http://www.nflrc.hawaii.edu/NetWorks/NW06/> (1997).
- [46] Sara Kiesler. 2005. Fostering common ground in human-robot interaction. In *Robot and Human Interactive Communication, 2005. ROMAN 2005. IEEE International Workshop on*. IEEE, 729–734.
- [47] Young Ji Kim, David Engel, Anita Williams Woolley, Jeffrey Yu-Ting Lin, Naomi McArthur, and Thomas W Malone. 2017. What makes a strong team?: Using collective intelligence to predict team performance in League of Legends. In *CSCW*. 2316–2329.
- [48] Ross A Knepper, Christoforos I Mavrogiannis, Julia Proft, and Claire Liang. 2017. Implicit communication in a joint action. In *Proceedings of the 2017 ACM/IEEE International Conference on Human-Robot Interaction*. ACM, 283–292.
- [49] Ross A Knepper, Stefanie Tellex, Adrian Li, Nicholas Roy, and Daniela Rus. 2015. Recovering from failure by asking for help. *Autonomous Robots* 39, 3 (2015), 347–362.
- [50] Hector J Levesque. 1989. A knowledge-level account of abduction. In *Proceedings of the 11th International Joint Conference on Artificial Intelligence*, Vol. 2. Morgan Kaufmann Publishers Inc., 1061–1067.
- [51] Stephen C Levinson. 2000. *Presumptive meanings: The theory of generalized conversational implicature*. MIT Press.
- [52] Claire Liang, Julia Proft, and Ross A Knepper. 2017. Implicature-based inference for socially-fluent robotic teammates. (2017).
- [53] Victor Ei-Wen Lo and Paul A Green. 2013. Development and evaluation of automotive speech interfaces: useful information from the human factors and the related literature. *International Journal of Vehicular Technology* 2013 (2013).
- [54] Derek Lomas, Kishan Patel, Jodi L Forlizzi, and Kenneth R Koedinger. 2013. Optimizing challenge in an educational game using large-scale design experiments. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, 89–98.
- [55] J Derek Lomas, Kenneth Koedinger, Nirmal Patel, Sharan Shodhan, Nikhil Poonwala, and Jodi L Forlizzi. 2017. Is difficulty overrated?: The effects of choice, novelty and suspense on intrinsic motivation in educational games. In *Proceedings of the 2017 CHI Conference on Human Factors in Computing Systems*. ACM, 1028–1039.
- [56] Christoforos I Mavrogiannis, Valts Blukis, and Ross A Knepper. 2017. Socially competent navigation planning by deep learning of multi-agent path topologies. In *Intelligent Robots and Systems (IROS), 2017 IEEE/RSJ International Conference on*. IEEE, 6817–6824.
- [57] Randolph A Miller, Harry E Pople Jr, and Jack D Myers. 1982. INTERNIST-I, an experimental computer-based diagnostic consultant for general internal medicine. *New England Journal of Medicine* 307, 8 (1982), 468–476.
- [58] Charles G Morgan. 1971. Hypothesis generation by machine. *Artificial Intelligence* 2, 2 (1971), 179–187.
- [59] Stefanos Nikolaidis, David Hsu, and Siddhartha Srinivasa. 2017. Human-robot mutual adaptation in collaborative tasks: Models and experiments. *The International Journal of Robotics Research* (2017), 0278364917690593.
- [60] Hirotaka Osawa. 2015. Solving Hanabi: Estimating hands by opponent's actions in cooperative game with incomplete information. In *Workshops at the Twenty-Ninth AAAI Conference on Artificial Intelligence*.
- [61] Lynne E Parker. 2000. Current state of the art in distributed autonomous mobile robotics. In *Distributed Autonomous Robotic Systems* 4. Springer, 3–12.
- [62] Charles S. Peirce. 1901. The proper treatment of hypotheses: A preliminary chapter, toward an examination of Hume's argument against miracles, in its logic and in its history. MS [R] 692. (1901).
- [63] Yun Peng and James A Reggia. 1990. *Abductive inference models for diagnostic problem-solving*. Springer.
- [64] André Pereira, Rui Prada, and Ana Paiva. 2012. Socially present board game opponents. In *Advances in Computer Entertainment*. Springer, 101–116.
- [65] Harry E Pople. 1973. On the mechanization of abductive logic. In *Proceedings of the 3rd International Joint Conference on Artificial Intelligence*. Morgan Kaufmann Publishers Inc., 147–152.
- [66] Harry E Pople, Jack D Myers, and Randolph A Miller. 1975. DIALOG: A model of diagnostic logic for internal medicine. In *Proceedings of the 4th International Joint Conference on Artificial Intelligence*, Vol. 1. Morgan Kaufmann Publishers Inc., 848–855.
- [67] Christopher Potts. 2015. Presupposition and implicature. In *The Handbook of Contemporary Semantic Theory*, Shalom Lappin and Chris

- Fox (Eds.). John Wiley & Sons, Chapter 6.
- [68] Zahar Prasov and Joyce Y Chai. 2008. What's in a gaze?: The role of eye-gaze in reference resolution in multimodal conversational interfaces. In *Proceedings of the 13th International Conference on Intelligent User Interfaces*. ACM, 20–29.
 - [69] François Recanatì. 2001. What is said. *Synthese* 128, 1 (2001), 75–91.
 - [70] Melissa J Rogerson, Martin R Gibbs, and Wally Smith. 2017. What can we learn from eye tracking boardgame play?. In *Extended Abstracts Publication of the Annual Symposium on Computer-Human Interaction in Play*. ACM, 519–526.
 - [71] Melissa J Rogerson, Martin R Gibbs, and Wally Smith. 2018. Cooperating to compete: The mutuality of cooperation and competition in boardgame play. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. ACM, 193.
 - [72] Benjamin Russell. 2012. Probabilistic reasoning and the computation of scalar implicatures. *Unpublished doctoral dissertation, Brown University* (2012).
 - [73] Albrecht Schmidt. 2000. Implicit human computer interaction through context. *Personal Technologies* 4, 2-3 (2000), 191–199.
 - [74] Alessandra Sciutti, Ambra Bisio, Francesco Nori, Giorgio Metta, Luciano Fadiga, and Giulio Sandini. 2013. Robots can be perceived as goal-oriented agents. *Interaction Studies* 14, 3 (2013), 329–350.
 - [75] Alessandra Sciutti, Laura Patane, Francesco Nori, and Giulio Sandini. 2014. Understanding object weight from human and humanoid lifting actions. *IEEE Transactions on Autonomous Mental Development* 6, 2 (2014), 80–92.
 - [76] Yasaman S Sefidgar, Prerna Agarwal, and Maya Cakmak. 2017. Situated tangible robot programming. In *Proceedings of the 2017 ACM/IEEE International Conference on Human-Robot Interaction*. ACM, 473–482.
 - [77] Mei Si, Stacy C. Marsella, and David V. Pynadath. 2009. Modeling appraisal in theory of mind reasoning. *Autonomous Agents and Multi-Agent Systems* 20, 1 (10 May 2009), 14.
 - [78] David Silver, Aja Huang, Chris J Maddison, Arthur Guez, Laurent Sifre, George Van Den Driessche, Julian Schrittwieser, Ioannis Antonoglou, Veda Panneershelvam, Marc Lanctot, et al. 2016. Mastering the game of Go with deep neural networks and tree search. *nature* 529, 7587 (2016), 484.
 - [79] David Silver and Joel Veness. 2010. Monte-Carlo planning in large POMDPs. In *Advances in Neural Information Processing Systems* 23, J. D. Lafferty, C. K. I. Williams, J. Shawe-Taylor, R. S. Zemel, and A. Culotta (Eds.). Curran Associates, Inc., 2164–2172. <http://papers.nips.cc/paper/4031-monte-carlo-planning-in-large-pomdps.pdf>
 - [80] Karin Slegers, Sanne Ruelens, Jorick Visser, and Pieter Duysburgh. 2015. Using game principles in UX research: A board game for eliciting future user needs. In *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems*. ACM, 1225–1228.
 - [81] Jack W Smith, John R Svirbely, Charles A Evans, Pat Strohm, John R Josephson, and Mike Tanner. 1985. RED: A red-cell antibody identification expert module. *Journal of Medical Systems* 9, 3 (1985), 121–138.
 - [82] Mark E Stickel. 1990. Rationale and methods for abductive reasoning in natural-language interpretation. In *Natural Language and Logic*. Springer, 233–252.
 - [83] Mark JH van den Bergh, Anne Hommelberg, Walter A Kusters, and Flora M Spieksma. 2016. Aspects of the cooperative card game Hanabi. In *Benelux Conference on Artificial Intelligence*. Springer, 93–105.
 - [84] Adam Vogel, Christopher Potts, and Dan Jurafsky. 2013. Implicatures and nested beliefs in approximate decentralized-POMDPs. In *ACL* (2). 74–80.
 - [85] Annika Wærn. 1994. Plan inference for a purpose. In *Proceedings of the Fourth International Conference on User Modeling*. 93–98.
 - [86] Joseph Walton-Rivers, Piers R. Williams, Richard Bartle, Diego Perez Liebana, and Simon M. Lucas. 2017. Evaluating and modelling Hanabi-playing agents. *CoRR* abs/1704.07069 (2017). <http://arxiv.org/abs/1704.07069>
 - [87] Deirdre Wilson and Dan Sperber. 1981. On Grice's theory of conversation. *Conversation and discourse* (1981), 155–78.
 - [88] Jason Wuertz, Sultan A Alharthi, William A Hamilton, Scott Bateman, Carl Gutwin, Anthony Tang, Zachary Toups, and Jessica Hammer. 2018. A design framework for awareness cues in distributed multiplayer games. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. ACM, 243.
 - [89] Taoshuai Zhang, Jie Liu, and Yuanchun Shi. 2012. Enhancing collaboration in tabletop board game. In *Proceedings of the 10th Asia Pacific Conference on Computer Human Interaction*. ACM, 7–10.
 - [90] Shengdong Zhao, Koichi Nakamura, Kentaro Ishii, and Takeo Igarashi. 2009. Magic cards: A paper tag interface for implicit robot control. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, 173–182.
 - [91] Wenjun Zhou, Wangcheng Yan, and Xi Zhang. 2017. Collaboration for success in crowdsourced innovation projects: Knowledge creation, team diversity, and tacit coordination. In *Proceedings of the 50th Hawaii International Conference on System Sciences*.