

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Concept Alignment as a Prerequisite for Value Alignment

Permalink

<https://escholarship.org/uc/item/9902x9hh>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 46(0)

Authors

Rane, Sunayana

Ho, Mark K

Sucholutsky, Ilia

et al.

Publication Date

2024

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Concept Alignment as a Prerequisite for Value Alignment

Sunayana Rane

Department of Computer Science
Princeton University
srane@princeton.edu

Ilia Sucholutsky

Department of Computer Science
Princeton University
is2961@princeton.edu

Mark Ho

Department of Computer Science
Stevens Institute of Technology
mark.ho@stevens.edu

Thomas L. Griffiths

Departments of Psychology & Computer Science
Princeton University
tomg@princeton.edu

Abstract

Value alignment is essential for building AI systems that can safely and reliably interact with people. However, what a person values—and is even capable of valuing—depends on the concepts that they are currently using to understand and evaluate what happens in the world. The dependence of values on concepts means that *concept alignment* is a prerequisite for value alignment—agents need to align their representation of a situation with that of humans in order to successfully align their values. Here, we formally analyze the concept alignment problem in the inverse reinforcement learning setting, show how neglecting concept alignment can lead to systematic value mis-alignment, and describe an approach that helps minimize such failure modes by jointly reasoning about a person’s concepts and values. Additionally, we report experimental results with human participants showing that humans reason about the concepts used by an agent when acting intentionally, in line with our joint reasoning model.

Introduction

People’s thoughts and actions are fundamentally shaped by the concepts they use to represent the world and formulate their goals. Imagine watching someone waiting to cross a busy intersection. Making sense of their behavior requires understanding their representation of things like “the crosswalk,” “the road,” “the bike lane,” and “the right of way.” For instance, it is important to take into account whether someone understands or is aware of the part of the street designated the “bike lane” while they wait since otherwise their intentions could be misinterpreted (e.g., a naïve observer might think someone standing in the bike lane is *trying* to get hit by a bicycle). Yet, current approaches to inferring human goals, rewards, and values (e.g., standard inverse reinforcement learning (Abbeel & Ng, 2004) and value alignment (Hadfield-Menell, Russell, Abbeel, & Dragan, 2016) largely neglect the possibility that a machine observer and a human actor can have misaligned concepts. Our goal in this work is to formally state the problem of *concept alignment*, begin to explore algorithmic solutions, and compare these solutions to human judgments.

To formalize concept alignment, we draw on the recently proposed framework of *value-guided construal* (Ho et al., 2022), which provides a computational account of how humans form simplified representations of problems in order to solve them. A *construal* is a particular interpretation of a problem in terms of a set of concepts and related causal affordances: for example, if one understands the concept of the bike lane and includes it in their current construal, they are

aware of the fact that bicycles are often on the bike lane, cars generally avoid the bike lane, you might get hit if you stand in the bike lane, etc. People often prefer simpler construals since they are less cognitively effortful (Ho, Cohen, & Griffiths, 2023), but this can affect the quality of one’s actions—e.g., if you fail to distinguish the bike lane from the sidewalk, you might stand in a place where a bicycle will hit you!

Work on inferring human preferences and values is often done in the framework of inverse-reinforcement learning (IRL) (Abbeel & Ng, 2004; Hadfield-Menell et al., 2016; Ho, Littman, MacGlashan, Cushman, & Austerweil, 2016) and inverse planning (Baker, Saxe, & Tenenbaum, 2009). However, a key property of virtually all existing IRL methods is that they assume behavior emerges from a planning process that produces optimal or noisy-optimal policies (Abbeel & Ng, 2004; Loftin et al., 2016; Ziebart, Maas, Bagnell, & Dey, 2008). This assumption is problematic because it is false (Simon, 1955; Tversky & Kahneman, 1974). An alternative perspective that has been developed over the past few years is that people are *resource-rational*—that is, they think and act rationally, but are subject to cognitive limitations on time, memory, or attention (Lieder & Griffiths, 2020). The key idea of value-guided construals is that people do not necessarily use all concepts available when representing a problem in order to make efficient use of limited attention (e.g., ignoring certain details of obstacles when navigating through a Grid-World). A major research challenge for IRL, value alignment, and cognitive science is incorporating these ideas into estimating human preferences and values (Ho & Griffiths, 2022; Evans, Stuhlmüller, & Goodman, 2016; Zhi-Xuan, Mann, Silver, Tenenbaum, & Mansinghka, 2020; Kwon et al., 2020; Alanqary et al., 2021; Chan, Critch, & Dragan, 2021; Laidlaw & Dragan, 2022).

We propose a theoretical framework for formally introducing concepts (in the form of construals) to inverse reinforcement learning and show that conceptual misalignment (i.e., failing to consider construals) can lead to severe value mis-alignment (i.e., reward mis-specification; large performance gap). We validate these theoretical results with a case study using a simple gridworld environment where we find that IRL agents that jointly model construals and reward outperform those that only model reward. Finally, we conduct a study with human participants and find that people do model construals, and that their inferences about rewards are a much

4000

closer match to the agent that jointly models construals and rewards. Our theoretical and empirical results suggest that the current paradigm of trying to directly infer human reward functions or preferences from demonstrations is insufficient for value-aligning AI systems that interact with real people; it is crucial to also model and align on the concepts people use to reason about the task in order to understand their true values and intentions.

Introducing Construals into Inverse Reinforcement Learning

We begin by reviewing the basic formalism for sequential decision-making before turning to construals and the inverse planning problem.

Markov decision processes (MDPs)

We represent sequential decision-making tasks as Markov decision-processes (MDPs) $M = \langle \mathcal{S}, \mathcal{A}, P_0, T, R, \gamma \rangle$, where $\mathcal{S}$ is a state space; $\mathcal{A}$ is an action space; $P_0 : \mathcal{S} \rightarrow [0, 1]$ is an initial state distribution; $T : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow [0, 1]$ is a transition function; $R : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ is a real-valued reward function; and $\gamma \in [0, 1]$ is a discount rate. A (stochastic) policy is a conditional probability distribution that maps states to distributions over actions, $\pi : \mathcal{S} \rightarrow \Delta(\mathcal{A})$. We denote the Markov chain resulting from following policy π on an MDP with dynamics T as $T^\pi(s' | s) = \sum_a \pi(a | s) T(s' | s, a)$.

We consider standard (unregularized) and entropy-regularized solutions to MDPs. In the unregularized setting, the value function associated with a policy π on an MDP with dynamics T and reward function R maps each state to the expected cumulative, discounted reward that results from following π : $V_{(R,T)}^\pi(s) = \sum_a \pi(a | s) [R(s, a) + \gamma \sum_{s'} T(s' | s, a) V_{(R,T)}^\pi(s')]$. The state occupancy function (also known as the successor representation) associated with a policy π on an MDP with dynamics T is the expected discounted visitations to a state s^+ starting from a state s , $\rho_T^\pi(s; s^+) = \mathbf{1}[s^+ = s] + \gamma \sum_{s'} T(s' | s) \rho_T^\pi(s'; s^+)$. The optimal value function for an MDP M maximizes value at each state, $V_{(R,T)}^*(s) = \max_a \{R(s, a) + \gamma \sum_{s'} T(s' | s, a) V_{(R,T)}^*(s')\}$.

In the entropy-regularized setting, the value of a policy π on MDP M is modified to include an entropy term, which penalizes action distributions that are more deterministic: $H(\pi(\cdot | s)) = -\sum_a \pi(a | s) \ln\{\pi(a | s)\}$. When this penalty is parameterized by a weight β , we denote the optimal entropy-weighted value function as $V_{(R,T)}^\beta(s) = \max_\pi \{\sum_a \pi(a) [R(s, a) + \sum_{s'} T(s' | s, a) V_{(R,T)}^\beta(s')] + \beta H(\pi)\}$.

Inverse Reinforcement Learning (IRL)

The standard IRL problem formulation involves an *observer* attempting to estimate the reward function of an expert *demonstrator* based on observed behavior (Abbeel & Ng, 2004). This can be formalized as Bayesian inference, where given a trajectory of expert acting in the task, $\zeta =$

$\{\langle s_0, a_0 \rangle, \langle s_1, a_1 \rangle, \dots, \langle s_T, a_T \rangle\}$, the observer infers the demonstrator’s reward function, R :

$$P(R | \zeta) = \frac{P(\zeta | R)P(R)}{P(\zeta)}. \quad (1)$$

To calculate the likelihood of a trajectory ζ given a reward function R , it is typically assumed that the observer has knowledge of the dynamics of the demonstrator’s task, T . Then, the likelihood is the probability of the trajectory being generated by the optimal policy under a candidate R :

$$P(\zeta | R) = \prod_{\langle s_t, a_t \rangle \in \zeta} \pi_{(R,T)}^\beta(a_t | s_t). \quad (2)$$

Inverse Construal

The inverse construal problem considers the possibility that although a resource-limited demonstrator is acting in a task with a particular dynamics T , they may not be *planning their actions* with respect to the fully-detailed dynamics. Rather, the demonstrator’s behavior results from planning with respect to a *construed task dynamics*, $\tilde{T}$, that is simpler or easier to solve.

Thus, an observer that takes into account the resource limitations faced by human planners should instead be aiming to solve an inference problem that incorporates the possibility of alternative task construals. Formally, this is the problem:

$$P(R, \tilde{T} | \zeta) = \frac{P(\zeta | R, \tilde{T})P(R, \tilde{T})}{P(\zeta)}, \quad (3)$$

where the likelihood is given by

$$P(\zeta | R, \tilde{T}) = \prod_{\langle s_t, a_t \rangle \in \zeta} \pi_{(R,\tilde{T})}^\beta(a_t | s_t). \quad (4)$$

Consequences of not considering construals

How bad can the estimate of R be when assuming the true dynamics T versus attempting to estimate the demonstrator’s construal $\tilde{T}$? If we use a maximum causal entropy formulation of IRL to get an estimated policy $\hat{\pi}^{\text{InvRL}}$ and compare this to the estimated policy assuming the demonstrator is using a construal, $\hat{\pi}^{\text{InvCon}}$, then the learner’s performance gap on the true task is (Viano, Huang, Kamalaruban, Weller, & Cevher, 2021):

$$|v_{(R,T)}^{\hat{\pi}^{\text{InvCon}}} - v_{(R,T)}^{\hat{\pi}^{\text{InvRL}}}| \leq \frac{\gamma \cdot |R|^{\max}}{(1 - \gamma)^2} \cdot \max_{s,a} \|T(\cdot | s, a) - \tilde{T}(\cdot | s, a)\|_1$$

where $|R|^{\max} = \max_{s,a} |R(s, a)|$. This is a tight bound, and thus the risk associated with not modeling construals (i.e., the potential size of the gap) grows rapidly when either the discounting, the maximum reward, or the task mismatch increases. In other words, if the observer has an inaccurate estimate of the transition function the actor uses to plan, they may drastically mis-estimate the reward function that motivated behavior. This provides a formal expression of our introductory example, in which failing to consider that a person does not know about or is unaware of a bike lane might lead one to interpret standing in the bike lane as indicating a *desire* to be hit by a bicycle.

A Simple Example of Concept Misalignment

Our theoretical analysis shows that concept misalignment creates risk of value misalignment (i.e., a large performance gap *can* exist). We now aim to show that the performance gap indeed exists in practice using a simple example.

First, we flesh out our bike lane example into a city navigation case study. Suppose Alice is trying to navigate a city to get a cup of coffee. There is a mom-and-pop bakery where she could get her favorite pastry and a delicious coffee, and a fast food franchise where she will have to wait in line and overpay but will still get a decent coffee. Given the choice, she would strongly prefer to go to the bakery. There are areas she can walk through (i.e., streets) and areas she cannot go through (i.e., locked buildings). However, there are also some unlocked buildings that she could cut through if only she noticed that they are unlocked. If she realizes both the bakery and the fast food place are accessible, she will always choose to go to the bakery (regardless of distance). However, if only the latter seems accessible, she will go get her coffee there. We visualize this setup in Figure 1. Now consider the case where the bakery is inside of a closed courtyard and the only way to reach it is to go through an unlocked building, but the fast food place is outside of the courtyard. If Alice is unaware that there are unlocked buildings that give access to the courtyard, she may end up going to the fast food place. An observer who does not take into account that Alice does not realize there are unlocked buildings to cut through would incorrectly infer that she prefers the fast food place (see Figure 1 third grid from the left).

To investigate the impact of modeling (or not modeling) a construal on value alignment between a human demonstrator and a machine IRL agent in scenarios such as this city navigation example, we use *blocks and notches* maze tasks similar to those developed by (Ho et al., 2023) to study rigidity in people’s construals (Figure 1). Each blocks-and-notches maze consists of a start state depicted as a blue circle (e.g., where Alice starts off), a high-value goal (e.g., the bakery) depicted as a pink or yellow square, a low-value goal (e.g., the fast food place) depicted as a yellow or pink square, and blue 3×3 blocks (e.g., buildings). The dark blue squares prevent movement (e.g., locked buildings), but the light blue notches (e.g., unlocked buildings) permit movement through the blocks. This environment allows us to simulate scenarios analogous to the city navigation ones described above.

Notches

In our simulations, notches (represented by light blue squares within the 3×3 blue blocks) are shortcuts through the grid. The idea is that while everyone is shown the same view when tasked with navigating the grid, only some demonstrators notice and learn how to use the notches; others ignore the light blue vs dark blue distinction and treat the entire 3×3 blue block as an obstacle, which they navigate around. In other words, people with different construals of the same ground truth grid learn different paths (Ho et al., 2023).

A standard IRL agent trying to infer a demonstrator’s rewards in this task is misaligned at the concept level because it assumes an optimal policy (and therefore has no notion that a human might not understand notches or how to use them). Humans, as we have discussed before, often act in ways that are not conventionally considered optimal or even rational. The IRL agent, without an understanding of the different construals people are using to understand the grid, draws incorrect conclusions about people’s values (rewards).

Of the four scenarios (Figure 1) used in our experiments, the two on the bottom depict routes taken by simulated demonstrators who did not realize they could walk through notches. The near (pink) goal is unreachable without using notches; think of it as the bakery enclosed on all sides by buildings (3×3 dark blue blocks) some of which are unlocked (light blue notches). The grids on the top show the trajectory of a simulated demonstrator who has learned that a notch is a shortcut, and has used the notch to form a more efficient path to their preferred goal. On the bottom are the trajectories of a simulated demonstrator who only knows the blue 3×3 blocks are obstacles, without paying attention to the fact that some sub-blocks (notches) are not obstacles at all. Looking at these trajectories on right, the IRL agent which does not have any notion of construals and assumes an optimal policy would naturally assume the human demonstrator has a value-related reason for avoiding the pink goal, and would thus assume that the yellow goal has a higher reward. Thus we see value misalignment emerge as a consequence of concept misalignment between the human demonstrator and the IRL agent.

Value misalignment

In our reinforcement learning framework, we use rewards as a proxy for values. To demonstrate how concept misalignment can lead to value misalignment between humans and machines, we employ an inverse reinforcement learning agent to infer the human’s values (reward function). Without knowledge of the construals (different understanding of notches), the agent might misattribute the path to a higher reward value for the chosen goal, not realizing the other goal may in fact have a higher reward, but may be impossible to reach without using/paying attention to notches.

As a measure of how alignment at the construal level can improve value alignment, we compare the posterior probability $P(R, \tilde{T} | \zeta)$ when jointly modeling the reward and the construal, to $P(R | \zeta)$, the standard IRL posterior which assumes that the trajectory is coming from a policy optimal with respect to the true transition function.

We run inference to calculate these probabilities for three GridWorlds shown in Figure 2 using both the reward-only model and the joint reward and construal model. The full twelve trajectories used for inference (four scenarios over each of the three GridWorlds) are shown in Figure 3.

Model results

We find that the joint reward and construal model performs on par with the reward-only model in the settings where the

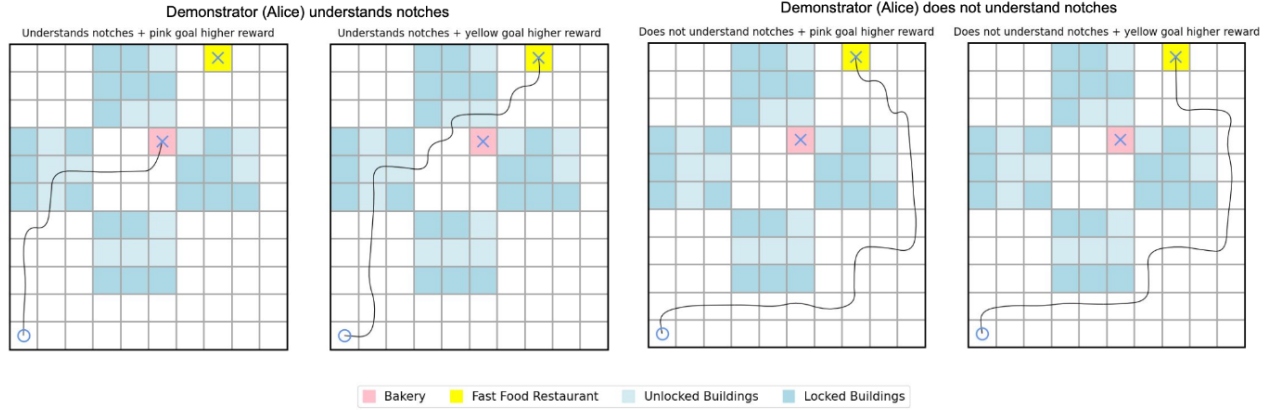


Figure 1: Four trajectories produced by different combinations of rewards and construals. The two trajectories on the right with the construal "Does not understand notches" look similar, because the preferred (pink) goal is impossible to reach when not construing notches.

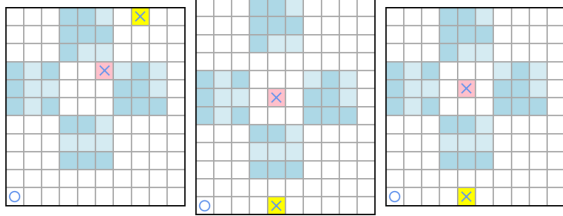


Figure 2: The three GridWorlds used in our experiments.

demonstrator "understands notches", but significantly outperforms the reward-only model in inferring values in the difficult "Does not understand notches" scenario where the preferred (pink) goal is inaccessible without using notches. In this scenario, the joint model correctly infers that the pink goal has a higher reward even though the demonstrator visits it in only one out of the three demonstrations for that scenario. The reward-only model fails completely, inferring confidently and incorrectly that the yellow goal has a higher reward due to the higher number of visits. See Figure 3 for a full comparison.

Testing the Model Predictions

We have now shown both theoretically and in simulation that a performance gap due to concept misalignment is not only possible but also plausible. But if vanilla inverse reinforcement learning is insufficient for inferring people's intent and preference in the real world, then how do people manage to do these things in practice? We now show that humans a) are highly adept at reasoning about construals, and b) use their knowledge of construals when making inferences about others' paths.

In this behavioral experiment, we gave 100 human participants the same four trajectories given to the two IRL agents

Table 1: Proportion of human participants who correctly inferred rewards for each scenario.

Scenario	Prop.
Understands notches + pink goal higher reward	0.99
Doesn't understand notches + pink goal higher reward	0.70
Understands notches + yellow goal higher reward	0.98
Doesn't understand notches + yellow goal higher reward	0.98

Note: All statistically significant at $p < .001$ using a one-sided binomial test against the null hypothesis of random chance.

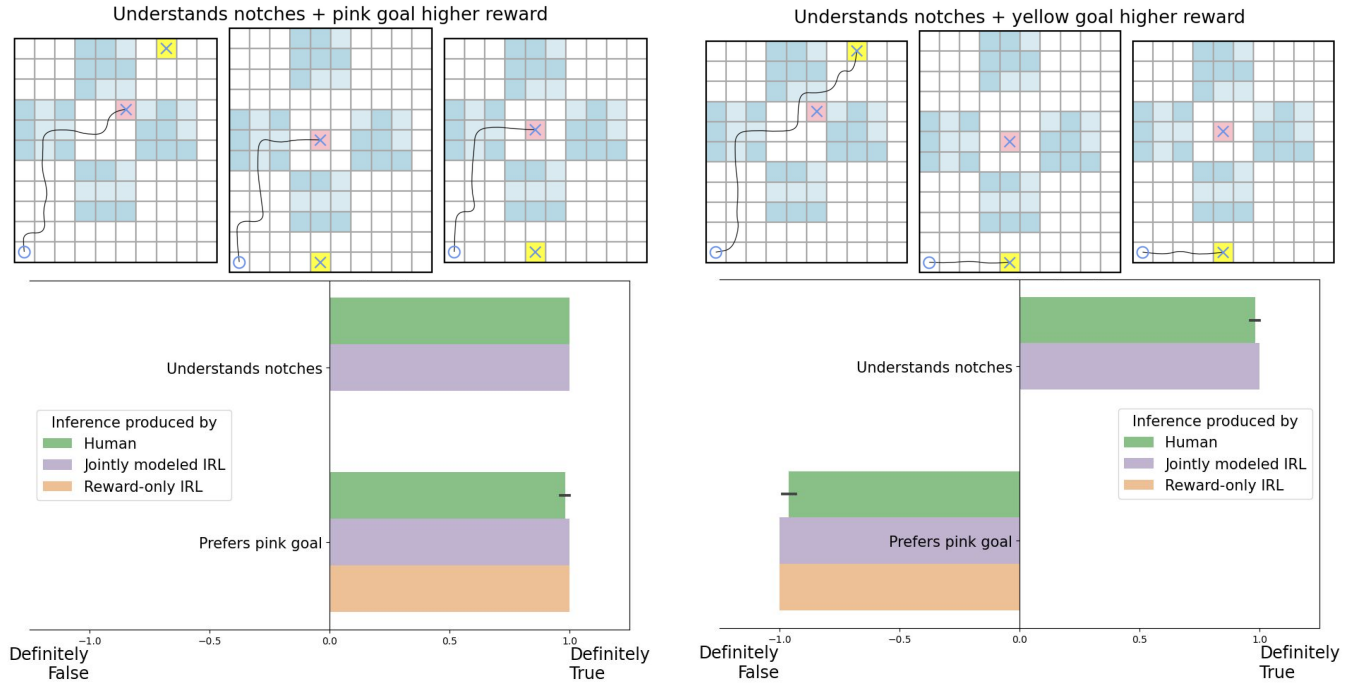
Table 2: Correlations between model and human inferences of reward

	Pearson correlation	p
Reward-only IRL	0.757	0.242
Jointly-modeled IRL	0.970	0.029

(Figure 1) for three different gridworlds, then asked them to make the same inferences. Each participant was shown a live replay of each trajectory, and then asked to infer (Figure 4) whether the person who took this route realized they could walk through notches, and if they preferred the pink goal to the yellow goal. Participants were asked to respond true or false to each question, and these responses were then mapped to scores of 1 or -1 when computing the results. These two questions map naturally to the posterior of the IRL algorithm's construal and reward inferences about each goal. We also scale the IRL posterior inferences to this -1 to 1 scale for direct comparison with the human judgments.¹

¹ Code and data: https://osf.io/hd9vc/?view_only=967b0c2f981d4a87bf4d21ff818f1322

Demonstrator understands notches



Demonstrator does not understand notches

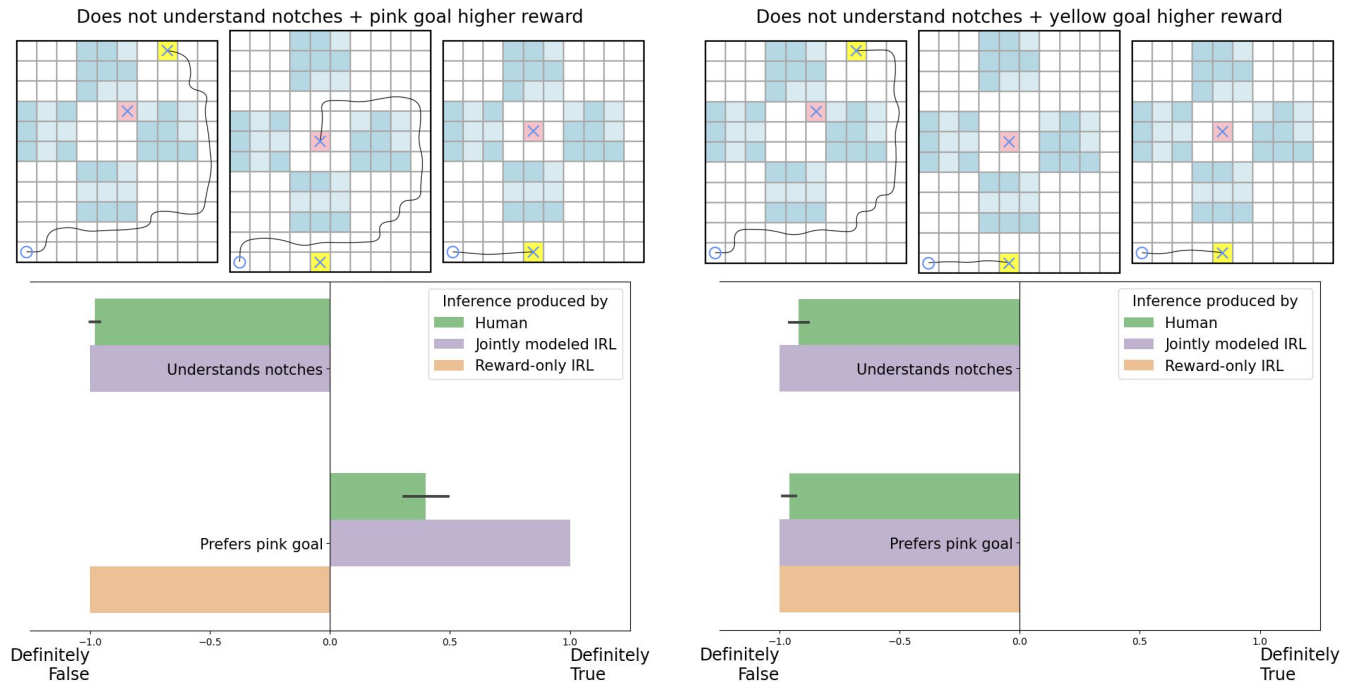


Figure 3: Inferences produced by humans and the two models. In the most difficult "Does not understand notches" scenario in the lower left corner, jointly modeling construals and rewards allows the IRL algorithm to successfully infer that the pink goal has higher reward, despite not being the most frequently accessed goal. Human subjects also make this inference. The reward-only IRL agent answers incorrectly and confidently that the yellow goal has higher reward.

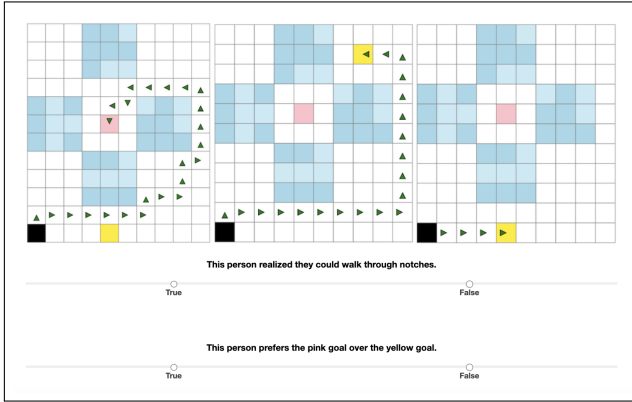


Figure 4: One frame of the data collection process. In this example, the person prefers the pink goal but does not realize they can walk through notches. An IRL algorithm that does not consider construals would assume that the person is intentionally choosing the yellow goal over the pink most of the time, and miscalculate the yellow goal having a higher reward, when in fact it is the other way around.

Results

There are three components to our results: human data, IRL inference when jointly modeling rewards and construals, and IRL inference when modeling only reward. These results are shown side-by-side in Figure 3. The posteriors of the IRL inference are scaled to match the -1 to 1 scale of the human data. Error bars for human data are one standard error from the mean over all 100 participants, for each question of each trajectory.

We show that humans completing the same inference task as the IRL agents successfully use construals to make more accurate reward inferences (see Table 1), producing judgments that are consistent with the joint reward and construal model (see Figure 3). In particular, there was a statistically significant difference in the number of people indicating the demonstrator preferred the pink goal in the conditions where the demonstrator did not understand notches ($\chi^2(1) = 40.05, p < .0001$), consistent with the predictions of the joint model and inconsistent with the reward-only model. We also calculated correlations between human reward inferences and model reward inferences, which demonstrate that the joint reward and construal model is highly correlated with human reward inferences in the same scenarios (see Table 2).

Discussion

We formulate the problem of concept alignment within the framework of value-guided construals. When people are faced with a task, they often do not represent it in full detail and instead engage in simplification strategies to make more efficient use of limited cognitive resources (Ho et al., 2022). As a result, people may use simplified concepts that lead to different behaviors than if they had represented the task in

complete detail. Our main goal here has been to formalize the inverse problem of estimating what simplified concepts people are using and show how such an approach is needed for successful value alignment and IRL in a simple setting.

In the most difficult scenario of the “Does not understand notches” construal, the reward-only IRL agent confidently makes a very incorrect inference (see Figure 3, because it does not model the construal. Modeling construals and allowing for alignment at a conceptual level enables the IRL algorithm to correctly infer human rewards and values instead of confidently making an incorrect inference. Modeling construals also brings the IRL behavior closer to the human participants’ behavior, because it creates a shared conceptual framework which enables more accurate reasoning about another person’s rewards and values.

Limitations and future work

While our theoretical results apply to many different settings with different types of construals, our experiments focused on a case study where we had knowledge of which features might form the basis for construal, and where we could control these features. In real-world settings, there are countless more features and a simple hard-coded construal model like the one we used for our experiments would clearly be insufficient. However, our goal was not to suggest that the specific simplified construal model we used for the experiments is the solution to the problem of modeling human intent and preferences, but rather to demonstrate that concept misalignment is, in fact, a problem that AI researchers need to focus on to make progress on value alignment. In future work, we intend to test this approach of using concept alignment to support value alignment in a wider variety of settings with human experts as well as develop inference algorithms that can scale to larger reward and construal spaces. More broadly, we hope that demonstrating the critical importance of concept alignment to the larger goal of value alignment will open the door to future work characterizing concept and value alignment in real-world settings.

Conclusion

Human decision-making is complex, context-driven, and resource-dependent and we have to keep that in mind when we try to teach AI systems to infer human values. We have laid out a framework introducing the notion of concepts into inverse reinforcement learning and showed both theoretically and empirically that without concept alignment it may often be impossible to achieve value alignment. We also showed in a behavioral study that people reason about each other’s concepts when making inferences about each other’s goals and values. We hope that these results encourage other researchers to work on concept alignment as a crucial component in value alignment, effective human-AI interaction, and the development of safe autonomous agents.

Acknowledgements. This work was supported by grants from the Office of Naval Research (ONR grant number N00014-22-1-2813) and the NOMIS Foundation.

References

- Abbeel, P., & Ng, A. Y. (2004). Apprenticeship learning via inverse reinforcement learning. In *Proceedings of the twenty-first International Conference on Machine Learning*.
- Alanqary, A., Lin, G. Z., Le, J., Zhi-Xuan, T., Mansinghka, V. K., & Tenenbaum, J. B. (2021). Modeling the mistakes of boundedly rational agents within a Bayesian theory of mind. *arXiv preprint arXiv:2106.13249*.
- Baker, C. L., Saxe, R., & Tenenbaum, J. B. (2009). Action understanding as inverse planning. *Cognition*, 113(3), 329–349.
- Chan, L., Critch, A., & Dragan, A. (2021). Human irrationality: both bad and good for reward inference. *arXiv preprint arXiv:2111.06956*.
- Evans, O., Stuhlmüller, A., & Goodman, N. (2016). Learning the preferences of ignorant, inconsistent agents. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 30).
- Hadfield-Menell, D., Russell, S. J., Abbeel, P., & Dragan, A. (2016). Cooperative inverse reinforcement learning. *Advances in Neural Information Processing Systems*, 29.
- Ho, M. K., Abel, D., Correa, C. G., Littman, M. L., Cohen, J. D., & Griffiths, T. L. (2022). People construct simplified mental representations to plan. *Nature*, 606(7912), 129–136.
- Ho, M. K., Cohen, J. D., & Griffiths, T. (2023, Mar). *Rational simplification and rigidity in human planning*. PsyArXiv. Retrieved from psyarxiv.com/aqxws
- Ho, M. K., & Griffiths, T. L. (2022). Cognitive science as a source of forward and inverse models of human decisions for robotics and control. *Annual Review of Control, Robotics, and Autonomous Systems*, 5(1), 33-53.
- Ho, M. K., Littman, M., MacGlashan, J., Cushman, F., & Austerweil, J. L. (2016). Showing versus doing: Teaching by demonstration. *Advances in Neural Information Processing Systems*, 29.
- Kwon, M., Biyik, E., Talati, A., Bhasin, K., Losey, D. P., & Sadigh, D. (2020). When humans aren't optimal: Robots that collaborate with risk-aware humans. In *Proceedings of the 2020 ACM/IEEE International Conference on Human-Robot Interaction* (pp. 43–52).
- Laidlaw, C., & Dragan, A. (2022). The Boltzmann policy distribution: Accounting for systematic suboptimality in human models. *arXiv preprint arXiv:2204.10759*.
- Lieder, F., & Griffiths, T. L. (2020). Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*, 43.
- Loftin, R., Peng, B., MacGlashan, J., Littman, M. L., Taylor, M. E., Huang, J., & Roberts, D. L. (2016). Learning behaviors via human-delivered discrete feedback: modeling implicit feedback strategies to speed up learning. *Autonomous Agents and Multi-Agent Systems*, 30(1), 30–59.
- Simon, H. A. (1955, February). A behavioral model of rational choice. *The Quarterly Journal of Economics*, 69(1), 99-118.
- Tversky, A., & Kahneman, D. (1974). Judgment under uncertainty: Heuristics and biases. *Science*, 185(4157), 1124-1131.
- Viano, L., Huang, Y.-T., Kamalaruban, P., Weller, A., & Cevher, V. (2021). Robust inverse reinforcement learning under transition dynamics mismatch. *Advances in Neural Information Processing Systems*, 34, 25917–25931.
- Zhi-Xuan, T., Mann, J., Silver, T., Tenenbaum, J., & Mansinghka, V. (2020). Online Bayesian goal inference for boundedly rational planning agents. *Advances in Neural Information Processing Systems*, 33, 19238–19250.
- Ziebart, B. D., Maas, A., Bagnell, J. A., & Dey, A. K. (2008). Maximum entropy inverse reinforcement learning. In *Proceedings of the 23rd National Conference on Artificial Intelligence* (p. 1433–1438). AAAI Press.