

Grounding Theory of Mind in Embodied AI

Hoa Thi Nguyen

Department of Applied Data Science, Institute for Energy Technology, Norway. E-mail: hoathi.nguyen@ife.no

Embodied AI studies agents that learn, perceive, and act through physical interaction with the world. Learning from human demonstrations is especially promising given the abundance of observable human behavior, but transferring knowledge from humans to robots remains challenging due to differences in embodiment, capability, and perspective. Inverse reinforcement learning (IRL) provides a mechanism for inferring goals from behavior, yet it often conflates true intentions with patterns shaped by physical or cognitive constraints. Human actions reflect not only what people want to achieve but also what is feasible or comfortable for them to do. We introduce PG-ToM, a perspective that treats physical and cognitive limitations as causal factors in action generation and integrates this view with Theory of Mind (ToM). Rather than attributing all behavioral regularities to preferences, PG-ToM frames actions as arising from the joint influence of goals and embodiment, allowing a learner to conceptually separate task objectives from feasibility-driven adaptations. We pose two research questions: (1) How does embodiment shape ToM in artificial agents? (2) How can ToM improve embodied AI's learning and adaptation in dynamic, interactive settings? Using a human–robot dishwashing scenario, we discuss how embodiment grounds ToM and how ToM enhances coordination, learning efficiency, and generalization. While the discussion centers on a low-stakes household task, the ideas extend to broader settings where interpreting human behavior in context is critical for effective human–robot interaction.

Keywords: Artificial intelligence, embodied AI, theory of mind

1. Introduction

In natural systems, intelligence develops in close coupling with embodiment: organisms evolve physical forms that shape how they perceive, move, and survive, and their intelligence emerges through ongoing interaction with the world. Intelligence is therefore not purely abstract computation but situated, embodied adaptation Pfeifer and Bongard (2006). In contrast, much of traditional AI has been developed in disembodied digital settings. Although robotics brings AI into physical contexts, transferring algorithms from simulation to reality remains difficult due to the “reality gap,” where real-world interaction introduces constraints and noise absent in virtual environments Salvato et al. (2021). As a result, systems that succeed in abstract domains may struggle in practice, motivating learning approaches that explicitly account for embodiment.

Human intelligence is also deeply social. Knowledge is shaped through interaction and shared experience Merleau-Ponty (2012), and humans reason not only about the physical world

but about each other. Theory of Mind (ToM), the ability to attribute beliefs, desires, and intentions to others Premack and Woodruff (1978), enables observers to interpret actions as goal-directed. However, human behavior is not purely rational; it reflects a mix of goals, preferences, habits, heuristics, and situational constraints Wood and Neal (2007); Tversky and Kahneman (1981). The same goal can therefore produce different behaviors across individuals or contexts. Interpreting behavior requires sensitivity to both intentions and the cognitive and contextual factors shaping their pursuit. Learning from human behavior is thus promising for embodied AI, given the abundance of demonstrations in daily life and online media. Yet many approaches focus on pattern extraction while overlooking intention and context. Machine ToM remains an emerging area without unified frameworks or evaluation standards Deshpande and Magerko (2024); Mao et al. (2024).

In this work, we explore how ToM can be integrated with embodied learning. We introduce a perspective called Physical Grounded Theory of Mind (PG-ToM), which treats physical and

cognitive constraints as causal factors in action generation. Rather than assuming that all regularities in behavior reflect preferences, PG-ToM views actions as shaped jointly by goals and embodiment. This distinction is especially important when learning from humans whose actions reflect feasibility, effort, and bounded rationality in addition to intention.

Using a robotic dishwasher scenario as a running example, we develop a conceptual framework for how ToM can support embodied learning. The dishwasher task provides a simple but illustrative setting in which human actions reflect both goals, such as safe and efficient placement, and constraints, such as reachability or dexterity.

This paper is guided by two research questions:

- (1) How does embodiment shape ToM in artificial agents?
- (2) How can ToM improve embodied AI's learning and adaptation in dynamic, interactive settings?

By examining these questions, we aim to clarify the role of ToM in embodied AI and to motivate learning approaches that interpret human behavior in a more context-sensitive and human-aware manner.

The rest of the paper is structured as follows. Section 2 presents the necessary background. Section 3 introduces the central ideas behind PG-ToM, and Section 4 illustrates how PG-ToM manifests in a human-robot dishwashing scenario. Section 5 revisits the research questions and discusses implications for reliability, interpretability, and safety. Finally, Section 6 concludes the paper.

2. Background

The concept of Theory of Mind (ToM) emerged in cognitive science to describe the capacity to attribute mental states to oneself and others, and to recognize that others may hold beliefs different from one's own Premack and Woodruff (1978). In computational treatments, agents' mental states are often represented using the Belief-Desire-Intention (BDI) model, which builds on Bratman's theory of practical reasoning Bratman (1987) and was later adopted in AI and multi-

agent systems as a software-oriented model of rational agency Rao and Georgeff (1995).

Embodied AI studies agents whose cognition and behavior arise from the tight coupling between control systems, physical bodies, and environments rather than from disembodied computation alone Pfeifer and Bongard (2006). This perspective introduces a central mismatch between what is directly observed—e.g., motor trajectories—and what must be inferred—e.g., intentions under physical and contextual constraints. Foundational computational ToM formulations address this gap through inverse planning, where observers infer latent goals or beliefs by inverting a generative model that assumes approximately rational action Mao et al. (2024). Inverse reinforcement learning (IRL) adopts a closely related premise, seeking reward functions that rationalize observed behavior Ng and Russell (2000). Within this framework, physical and social-cognitive influences are typically subsumed into a unified decision-theoretic model: states encode situations, uncertainty captures limited knowledge, and rewards summarize preferences Arora and Doshi (2021). While this abstraction is flexible, it rarely foregrounds embodiment or mental-state reasoning as distinct explanatory factors.

Bayesian action-understanding models partially address this limitation by modeling “rationality under constraints,” treating actions as probabilistic choices shaped by cost functions and environmental feasibility Baker et al. (2005). In robotics, the assumption of approximate rationality underlies noisy or Boltzmann-rational goal inference and many IRL pipelines, effectively acting as a ToM-inspired prior in which goals are revealed only through unknown costs Tabrez et al. (2020). BDI architectures have further been formalized as Bayesian ToM models that cast mental-state attribution as probabilistic inference Baker et al. (2009). Beyond prediction, ToM is also viewed as a prerequisite for social learning Bandura (1977), since it enables the Self-Other distinction—separating the demonstrator's perspective from the learner's own state. This distinction is essential in Learning from Observation (LfO) Torabi et al. (2019), where agents must reinter-

pret demonstrations within their own embodiment rather than directly replicate observed motions.

Recent evidence suggests that physical constraints are not merely noise but informative signals. Human movement reflects systematic trade-offs among task success, effort, comfort, and morphology; inverse optimal control studies repeatedly recover composite costs related to torque, work, and smoothness, indicating that deviations from task-optimality may be embodiment-optimal Berret et al. (2011). In collaborative robotics, incorporating kinematic feasibility and reachability constraints improves both prediction and assistance, as constraint-aware intent estimation accounts for the capabilities of all partners in physical interaction Shao et al. (2024). Assistance systems likewise achieve greater reliability by inferring goals from partial or noisy signals rather than cloning trajectories, reinforcing the value of intention-level inference Aronson and Admoni (2022).

Despite these advances, many robotic ToM applications still assume approximate optimality in task space, with embodiment represented as fixed penalties or ignored. Surveys note that noisyrational models typically presume a stable underlying reward and rarely capture time-varying physical limits such as fatigue, strength, or joint capacity Tabrez et al. (2020). Emerging human digital twin and fatigue-estimation models offer ways to represent these latent physiological states, yet they remain largely uncoupled from ToM-style inverse planning in closed-loop learning You et al. (2024).

3. Physically Grounded Theory of Mind

ToM improves robotic reliability by supporting intention-level inference and reducing ambiguity in human behavior interpretation, yet conventional ToM models often overlook the physical constraints shaping human actions. To address this gap, we introduce a physically grounded ToM (PG-ToM) framework that integrates embodiment factors into mental-state reasoning. PG-ToM rests on three core assumptions:

(1) Human actions are not purely goal-optimal:

They frequently reflect trade-offs between goals and physical feasibility.

- (2) Embodiment introduces systematic bias: Factors like fatigue, strength limitations, joint mobility, and reachability systematically influence task execution.
- (3) Reliable learning requires intention-level inference: AI systems should infer the intended outcome behind an action rather than merely reproducing the observed motor pattern.

To overcome these issues, PG-ToM explicitly treats embodiment as a causal component in action generation.

4. Experimental Scenario: Cobot Dishwashing

We consider a scenario in which a human and a collaborative robot (cobot) jointly load a dishwasher. The interaction operates in three modes:

- (1) Demonstration–Observation Mode: When the human is near the rack, the robot observes and internalizes human actions into its policy.
- (2) Deployment–Assistance Mode: When the human steps back, the robot executes its learned policy to assist with loading.
- (3) Human Corrective Feedback: If the robot behaves undesirably, such as placing fragile items unsafely, the human can intervene, providing corrective feedback that triggers policy updates.

Rather than replicating low-level motor actions, the cobot infers higher-level goals and preferences through inverse reinforcement learning (IRL). The emphasis is on goal inference rather than motion imitation.

The human’s primary objective is to load all items into the dishwasher. From this, the system infers context-dependent subgoals, such as maintaining rack stability, promoting drying efficiency through appropriate orientations, and ensuring safe handling of fragile items. Human demonstrators are treated as near-optimal, while allowing variability due to physical and cognitive differences.

Within PG-ToM, the focus shifts from trajectory imitation to alignment with the human’s be-

lief–desire–intention (BDI) structure. The aim is to recover underlying goals without overfitting to surface-level behavior.

4.1. Embodied Observation

PG-ToM relies on perception to capture both task context and embodiment cues. Multimodal sensing, such as RGB-D vision, supports:

- Object-centric perception: detecting dishes, cups, and utensils, estimating poses and orientations, and inferring attributes like category or fragility.
- Human activity tracking: pose estimation, hand tracking, and coarse indicators of cognitive or physical state.

In the framework of PG-ToM, beyond capturing task-relevant actions, the system also tracks indicators of physical constraints, such as posture adjustments, overextended reaching, repeated repositioning, or signs of exertion. The goal is not to construct a detailed biomechanical model, but to detect situations in which behavior is likely influenced by bodily limitations. Actions that seem indirect or inefficient are therefore marked as potentially constraint-driven.

4.2. Intention Inference

Classical IRL infer the latent reward functions from human behavior with the assumption that human actions are optimal to achieve the maximum award. Physical and cognitive constraints are not modeled in IRL as separate variables. Instead, they are incorporated implicitly through the modeling assumptions about the expert and environment. They influence IRL in three main ways:

- The transition function captures how placement actions change the state of the dishwasher (e.g., which slots are occupied, object orientations, rack balance). Human physical factors such as reachability limits, preferred grasp motions, or movement smoothness influence the demonstrations only insofar as they shape these observed state transitions. For example, if a human consistently loads

plates from front to back because of reach comfort, this pattern appears as a regularity in the transitions rather than as an explicit “reachability constraint.” As a result, IRL’s accuracy depends on how well these human–environment interaction dynamics are modeled.

- The inferred reward is typically represented as a linear combination of features that encode task-relevant preferences. In the dishwasher case, these features might capture plate alignment stability, drying-oriented orientations, separation of fragile items, or efficient space usage. Human factors such as fatigue, effort, or comfort are not inherently represented; they influence learning only if the state or feature design encodes proxies for them (e.g., penalizing long reach distances or awkward placements).
- If the human demonstrator does not have perfect information about the environment, for example being uncertain whether an item is clean, whether a slot is partially occupied, or whether an object is fragile, this uncertainty can be represented as a belief, meaning a probability distribution over possible states. The demonstrator’s actions are then interpreted as approximately optimal given this belief, rather than with respect to the true state of the dishwasher.

$$a_t = \arg \max_a Q^*(s_t, a),$$

where $Q^*(s_t, a)$ denotes the optimal action-value function, i.e., the expected cumulative reward obtained after taking action a in state s_t and thereafter following the optimal policy.

To disentangle human goals from embodiment effects, in PG-ToM, we model human behavior as arising not solely from reward optimization but from the interaction between goals, beliefs, intentions, and physical feasibility. Rather than attributing all systematic action patterns to the reward function, we augment the decision-making state to include internal and embodied factors such as belief states, intention or subgoal commitments, and physical state variables capturing fatigue, strength

limitations, or reachability.

$$x_t = (s_t, b_t, i_t, p_t)$$

s_t : external world state

b_t : belief state

i_t : intention or subgoal

p_t : physical state.

$$a_t = \arg \max_{a \in \mathcal{A}(x_t)} Q^*(x_t, a)$$

where $Q^*(x_t, a)$ is the optimal action-value function defined over the augmented state x_t , and $\mathcal{A}(x_t)$ denotes the feasible action set given embodiment constraints such as reachable placements only.

In the dishwasher setting, that sentence would mean that the robot does not treat every consistent human habit as a preference baked into the reward. Instead, it represents those habits as consequences of internal state and embodiment that change what the human knows, intends, and can physically do at each step.

The human belief state b_t captures the demonstrator’s subjective understanding and expectations about the task and environment. Beyond momentary perceptual uncertainty, it can include prior knowledge, habits, and assumptions formed through experience. In the dishwasher context, this may encompass beliefs about which items are fragile or valuable, expectations about how well certain orientations promote drying, perceived cleanliness, or assumptions about what “counts” as properly loaded. It may also reflect individual tendencies such as risk aversion (e.g., overestimating breakage risk) or confidence in one’s own placements. Thus, the belief state encodes how the human interprets the situation, not just what is objectively true.

The intention/commitment state i_t represents the demonstrator’s active goals, priorities, and task strategies at a given moment. Rather than a single fixed objective, humans dynamically shift between subgoals and modes of operation. This state can capture personal preferences and personality-linked tendencies, such as prioritizing neatness versus speed, being meticulous versus

satisficing, or valuing symmetry and visual order. It may also encode social or normative motivations (e.g., loading “the way my household expects”). Commitment strength can vary: a person may persist strongly with a preferred arrangement style or flexibly adapt when space is limited. In this sense, i_t reflects not only what the human wants to achieve, but how they prefer to achieve it.

The embodied state p_t represents the human’s physical capabilities, comfort, and momentary condition, which systematically shape feasible and preferred actions. This includes fatigue, strength, dexterity, joint mobility, and reachability, but can be broadened to stable traits such as height, handedness, motor habits, or age-related factors. It may also reflect comfort preferences (e.g., avoiding awkward bends), movement style (careful vs brisk), and personal tolerance for effort. These factors introduce consistent patterns in task execution that are not necessarily preferences over outcomes but constraints and biases arising from embodiment. Modeling p_t allows the system to attribute regularities in behavior to physical feasibility and comfort rather than misattributing them to reward preferences.

The transition dynamics then encode how these internal and physical states evolve over time and constrain the set of feasible actions.

The reward function encodes task-level objectives, such as completing the loading task safely and efficiently, while departures from apparent goal-optimality are attributed to bounded rationality and embodiment constraints. This separation enables inverse reinforcement learning to infer rewards that better reflect underlying human preferences, while systematic effects of physical and cognitive limitations are modeled through the dynamics rather than absorbed into the reward.

PG-ToM treats physical limitations as a causal factor in action generation. Rather than attributing observed placement patterns solely to latent preferences, PG-ToM interprets them as arising from the joint influence of task goals and physical feasibility. In this view, dish placements are understood not only as reflections of what an individual aims to achieve, but also as adaptations to their

momentary physical capabilities and constraints. This perspective allows the model to distinguish between goal-driven behavior and embodiment-driven adjustments, leading to more faithful inference of underlying objectives.

$$U(x_t, a) = Q_R(x_t, a) - \lambda C(p_t, s_t, a),$$

where $Q_R(x_t, a)$ denotes the reward-based action-value term, $C(p_t, s_t, a)$ is a cost function dependent on physical and environmental factors, and λ is a trade-off parameter balancing reward and cost.

$$a_t = \arg \max_a U(x_t, a).$$

PG-ToM is therefore to recover latent goals from demonstrations while factoring out embodiment-driven distortions.

4.3. Cross-Embodiment Translation

After inferring the task reward through the PG-ToM framework, the robot learns and plans within its own action space, consistent with its kinematics, reachability, and manipulation capabilities. In the dishwasher setting, this means the robot optimizes how it places dishes based on its own gripper design, joint limits, and collision constraints rather than imitating human motions directly. To improve sample efficiency, human demonstrations can be passed through a biomechanical-to-robot mapping that converts feasible human placements and orientations into robot-executable actions when possible, allowing the robot to reuse relevant aspects of the demonstrations. Crucially, the inferred goals, such as safe placement, rack stability, and drying-oriented orientations, remain fixed, while the execution strategy is adapted to the robot's embodiment.

5. Discussions

5.1. Addressing the Research Questions

Embodiment shapes ToM by constraining how actions can be generated and therefore how they should be interpreted. PG-ToM starts from the premise that behavior is not driven by goals alone, but by the interaction between goals and physical feasibility. When an agent observes a person

acting, it should not immediately treat recurring patterns as preferences. Instead, PG-ToM treats physical limitations such as reachability, dexterity, or fatigue as causal factors that influence what actions are likely at a given moment. This perspective reframes ToM from asking only "What does this person want?" to also asking "What can this person reasonably do right now?" In PG-ToM, embodiment provides part of the explanation for behavior, allowing the agent to avoid conflating physical convenience with true intention. As a result, inferred goals are more stable and more representative of what the person is actually trying to achieve.

PG-ToM illustrates how ToM can directly support learning and adaptation by separating goals from embodiment-driven execution. When a robot learns from demonstration, it often faces variability across users and situations. PG-ToM uses ToM-inspired reasoning to explain this variability partly through constraints rather than preferences. This allows the system to extract goal-level knowledge that generalizes across people while still adapting its execution to its own body and to the current partner. In dynamic interactions, PG-ToM also supports better anticipation and coordination. By reasoning about a person's likely goals while accounting for their physical situation, a robot can choose actions that are complementary rather than disruptive. For example, it can infer when a person avoids certain placements due to reach limits and assist accordingly. In this way, ToM is not only a tool for interpreting behavior but also a mechanism for adaptive, human-aware collaboration. More broadly, PG-ToM shows that ToM and embodiment should not be treated separately. ToM becomes more accurate when grounded in embodiment, and embodied AI becomes more adaptive when guided by ToM. Together, they enable systems that learn from humans in a way that is both more robust and more aligned with real-world behavior.

5.2. Implications for Reliability, Interpretability, and Safety

By separating task goals from embodiment constraints, PG-ToM reduces a common failure mode

in learning from demonstration where physical difficulty is misinterpreted as preference. This can improve reliability because the learned objectives are less sensitive to incidental factors such as human reachability or fatigue. In the dishwasher scenario, for example, the system is less likely to conclude that a rarely used rack position is undesirable when it may simply be hard to reach for a particular demonstrator. This separation can also support better generalization across users and embodiments. Reliability depends on how well physical and cognitive constraints are modeled. If the system incorrectly estimates what is feasible or comfortable for the human, it may misattribute behavior to goals or constraints. In addition, introducing more latent factors increases modeling complexity and may require more data to learn robustly.

PG-ToM promotes interpretability by offering a structured explanation of behavior. Actions can be explained in terms of goals, such as safe placement or drying efficiency, versus constraints, such as reachability or dexterity. This decomposition aligns well with how humans naturally reason about behavior and can make system decisions easier to audit and communicate. The richer internal representation can also introduce ambiguity. Multiple explanations may fit the same behavior, for example a careful placement could reflect fragility concerns or low confidence. Without additional signals, distinguishing between these explanations can be challenging.

Even in a non-critical domain like dishwasher loading, safety-relevant patterns emerge, such as handling fragile items or avoiding unstable configurations. By explicitly modeling constraints, PG-ToM can encourage conservative and physically feasible actions. When transferred to safety-critical domains, this perspective can help avoid over-aggressive plans that assume idealized human or robot capabilities. The model does not guarantee safety by itself. If the inferred goals or constraint models are wrong, the system may still produce unsafe behavior. Therefore, PG-ToM should be combined with explicit safety checks, physical limits, and task-level safeguards, especially in high-stakes applications.

Although the dishwasher task is low risk, the underlying idea of disentangling goals from embodiment can translate to more critical settings such as assistive robotics, healthcare, or driving. In these contexts, distinguishing what a person intends from what they are physically able to do becomes even more important. At the same time, higher-stakes applications demand stronger validation, clearer uncertainty handling, and integration with formal safety mechanisms.

6. Conclusion

This work presented PG-ToM, a framework that treats physical and cognitive constraints as causal contributors to action generation rather than conflating them with preferences. By separating task goals from embodiment-driven influences, PG-ToM enables more faithful inference of what a demonstrator is trying to achieve. In the dishwasher loading scenario, this allows a robot to distinguish between goal-relevant regularities, such as safe placement, stability, and drying-oriented orientations, and patterns that arise from reachability, fatigue, or individual execution style. The resulting goal representations are therefore more stable across users and more transferable to the robot's own embodiment.

Although we evaluated PG-ToM in a low-stakes household task, the underlying insight extends beyond dishwashing. Human behavior in many domains reflects a mixture of intention and constraint, and learning systems benefit from recognizing this distinction. PG-ToM highlights how Theory of Mind and embodiment can be integrated to produce more context-sensitive interpretations of behavior. Future work can explore richer signals, longer-term interactions, and applications in domains where robustness and safety are increasingly important, further advancing how embodied agents learn from and collaborate with people.

References

- Aronson, R. M. and H. Admoni (2022, June). Gaze complements control input for goal prediction during assisted teleoperation. In *Pro-*

- ceedings of Robotics: Science and Systems (RSS '22)*.
- Arora, S. and P. Doshi (2021). A survey of inverse reinforcement learning: Challenges, methods and progress. *Artificial Intelligence* 297, 103500.
- Baker, C. L., R. Saxe, and J. B. Tenenbaum (2009). Action understanding as inverse planning. *Cognition* 113(3), 329–349. Reinforcement learning and higher cognition.
- Baker, C. L., J. B. Tenenbaum, and R. R. Saxe (2005). Bayesian models of human action understanding. In *Proceedings of the 19th International Conference on Neural Information Processing Systems, NIPS'05*, Cambridge, MA, USA, pp. 99–106. MIT Press.
- Bandura, A. (1977). *Social Learning Theory*. Englewood Cliffs, NJ: Prentice-Hall.
- Berret, B., E. Chiovetto, F. Nori, and T. Pozzo (2011). Evidence for Composite Cost Functions in Arm Movement Planning: An Inverse Optimal Control Approach. *PLOS Computational Biology* 7(10), e1002183.
- Bratman, M. E. (1987). *Intention, Plans, and Practical Reason*. Cambridge: MA: Harvard University Press.
- Deshpande, M. and B. Magerko (2024, May). Embracing Embodied Social Cognition in AI: Moving Away from Computational Theory of Mind. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems, CHI EA '24*, New York, NY, USA, pp. 1–7. Association for Computing Machinery.
- Mao, Y., S. Liu, Q. Ni, X. Lin, and L. He (2024). A review on machine theory of mind. *IEEE Transactions on Computational Social Systems* 11(6), 7114–7132.
- Merleau-Ponty, M. (2012). *Phenomenology of Perception*. London and New York: Routledge.
- Ng, A. Y. and S. J. Russell (2000). Algorithms for inverse reinforcement learning. In *Proceedings of the Seventeenth International Conference on Machine Learning, ICML '00*, San Francisco, CA, USA, pp. 663–670. Morgan Kaufmann Publishers Inc.
- Pfeifer, R. and J. C. Bongard (2006). *How the Body Shapes the Way We Think: A New View of Intelligence*. Cambridge, MA, USA: MIT Press.
- Premack, D. and G. Woodruff (1978). Does the chimpanzee have a theory of mind? *Behavioral and Brain Sciences* 1(4), 515–526.
- Rao, A. S. and M. P. Georgeff (1995). Bdi agents: From theory to practice. In *Proceedings of the First International Conference on Multiagent Systems (ICMAS'95)*, San Francisco, CA, USA, pp. 312–319. IEEE Computer Society / AAAI Press.
- Salvato, E., G. Fenu, E. Medvet, and F. A. Pellegrino (2021). Crossing the reality gap: A survey on sim-to-real transferability of robot controllers in reinforcement learning. *IEEE Access* 9, 153171–153187.
- Shao, Y. S., T. Li, S. Keyvanian, P. Chaudhari, V. Kumar, and N. Figueroa (2024). Constraint-aware intent estimation for dynamic human-robot object co-manipulation.
- Tabrez, A., M. B. Luebbbers, and B. Hayes (2020). A survey of mental modeling techniques in human–robot teaming. *Current Robotics Reports* 1(4), 259–267.
- Torabi, F., G. Warnell, and P. Stone (2019, August). Recent Advances in Imitation Learning from Observation. In *Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence*, Macao, China, pp. 6325–6331. International Joint Conferences on Artificial Intelligence Organization.
- Tversky, A. and D. Kahneman (1981). The framing of decisions and the psychology of choice. *Science* 211(4481), 453–458.
- Wood, W. and D. T. Neal (2007, oct). A new look at habits and the habit-goal interface. *Psychological Review* 114(4), 843–863.
- You, Y., Y. Liu, and Z. Ji (2024). Human digital twin for real-time physical fatigue estimation in human-robot collaboration. In *2024 IEEE International Conference on Industrial Technology (ICIT)*, pp. 1–6.